

Gravierende Mängel beim Kündigungs-Button

Seit dem 1. Juli 2022 müssen in Deutschland tätige Unternehmen den sogenannten Kündigungs-Button auf ihren Webseiten anbieten. Die Verbraucherzentrale Bayern hatte systematisch Webseiten überprüft und bei der Mehrheit davon erhebliche rechtliche Mängel aufgedeckt. Ein Großteil bewegte sich überdies im Graubereich, teilt die Verbraucherzentrale Bayern mit.

Die Verbraucherverbände hätten insgesamt 152 Unternehmen abgemahnt. Lediglich auf 273 von 840 überprüften Websites fanden sich gesetzeskonforme Kündigungs-Buttons, heißt es aus Bayern weiter.

349 Webseiten ließen den vorgeschriebene Kündigungs-Button ganz vermissen. In 65 Fällen war er auf der Website versteckt, in 38 Fällen trug er eine unzulässige Beschriftung. Überdies wurden 339 weitere Verstöße im Zusammenhang mit der Bestätigungsseite und dem finalen Bestätigungs-Button festgestellt, teilen die Verbraucherschützer mit. Es hätten zum Beispiel Pflichtangaben gefehlt, oder es habe unzulässige Beschriftungen gegeben. Letztere müssen ebenfalls bestimmten Formalien genügen.

Bis Anfang November 2022 zeigten sich 86 Unternehmen einsichtig und unterschrieben die geforderte Unterlassungserklärung. In drei Fällen erwirkten die Verbraucherschützer eine einstweilige Verfügung, in 17 Fällen haben sie ein Klageverfahren vorbereitet oder bereits ein solches eingereicht.

Sicherheitsforscher Sönke Huster über Lücken im WLAN-Stack des Linux-Kernels



„Es reicht, wenn du dein WLAN anhast“

Über die Luft gehackt werden, nur weil das WLAN eingeschaltet ist? Im August hat Sönke Huster Sicherheitslücken im WLAN-Stack des Linux-Kernels gefunden, die einen solchen Angriff theoretisch ermöglicht hätten. Seine Entdeckung zeigt, wie wichtig es ist, Software ausführlich zu testen.

Über die Luft gehackt werden, nur weil das WLAN eingeschaltet ist? Im August hat Sönke Huster Sicherheitslücken im WLAN-Stack des Linux-Kernels gefunden, die einen solchen Angriff theoretisch ermöglicht hätten. Seine Entdeckung zeigt, wie

wichtig es ist, Software ausführlich zu testen.

Von Kathrin Stoll

Sönke Huster ist wissenschaftlicher Mitarbeiter am Secure Mobile Networking Lab (SEEM00) der TU Darmstadt. Im August 2022 hat er fünf Sicherheitslücken im WLAN-Stack des Linux-Kernels entdeckt. Mittlerweile gibt es Patches. Wir haben mit ihm über den Fund, seine Methodik und den Disclosure-Prozess gesprochen.



Der Sicherheitsforscher Sönke Huster hat fünf Sicherheitslücken im Linux-Kernel gefunden. Wie er das gemacht hat, verrät er im Gespräch mit c't. *Josephine Franz*

c't: Wie kommt man darauf, im Linux-Kernel nach Sicherheitslücken zu suchen?

Sönke Huster: Ich habe dieses Jahr meine Masterthesis über Bluetooth-Fuzzing unter Linux geschrieben. Die Idee kam von meiner Masterarbeitsbetreuerin Dr. Jiska Classen. Im Bluetooth-Stack habe ich dann auch ein paar kleine Sicherheitslücken gefunden. Dann wurde ich wissenschaftlicher Mitarbeiter am Secure Mobile Networking Lab von Prof. Matthias Hollick und es lag nahe, es auf WLAN auszuweiten. Aus Angreifersicht sind WLAN und Bluetooth super interessant und auch irgendwie ähnlich. Wenn ich dich hacken will, ist es ja viel cooler, ich kann das durch die Luft aus dem Raum nebenan machen, ohne dass ich dafür erst physisch auf deinen Rechner zugreifen können muss, um zum Beispiel einen USB-Stick einzustecken. Beide Protokolle sind dafür prädestiniert.

c't: Du hast gleich fünf Lücken im Linux-Kernel gefunden. Wie bist du dabei vorgegangen?

Huster: Die Methode, die ich verwende, heißt Fuzzing. Sie wurde in den Achtzigerjahren von Barton Miller [Professor der Informatik in Madison, Wisconsin, Anm. d. Red.] entdeckt, der sich über eine Telefonleitung auf Holzmasten remote auf seinem Arbeitsrechner einloggte. Bei Gewitter wurde die Übertragung des Signals gestört und seine Eingaben kamen verzerrt an. Das führte dazu, dass Programme abstürzten oder sich anders verhielten als erwartet. So kam man dahinter, dass man zufällige Eingaben nutzen kann, um Bugs und Sicherheitslücken zu finden und das Fuzzing – auch Fuzz-Testing – war erfunden. Heute verwendet man dazu sogenannte Fuzzer. Das sind im Grunde Programme, die die Eingabeschnittstellen von Programmen, Betriebssystemen oder Netzwerken mit zufälligen Daten fluten.

Mit komplett zufälligen Eingaben arbeitet man heute aber nicht mehr. Man kann das Verfahren verfeinern und Eingaben benutzen, die nah an denen sind, die das Target – in diesem Fall eben Linux in meiner VM – erwartet. Um WLAN zu untersuchen, lasse ich den Fuzzer WLAN-Pakete mit kleinen Anomalien an das Linux-System in meiner virtuellen Maschine schicken, die er fortlaufend verändert. Dabei beobachtet und dokumentiert der

Fuzzer, welcher Code im Kernel zur Verarbeitung der mutierten WLAN-Pakete getriggert wird. Man könnte auch sagen: welchen Weg ein Paket bei der Verarbeitung nimmt. Immer, wenn bei der Verarbeitung eines Pakets Code abgedeckt wurde, der vorher noch nicht ausgeführt wurde, nimmt der Fuzzer dieses Paket in sein Eingabeset auf und nutzt es als Ausgangspunkt für neue Mutationen. Diese veränderten Pakete schickt er dann wieder an den Kernel. Das Ganze passiert ein paar Tausend Mal pro Sekunde. Das Ziel ist es, möglichst viel Code „zu covern“, also durch die mutierten Eingaben Teile des Kernel-Codes abzudecken, die der Fuzzer noch nicht kennt. Coverage-Guided Mutational Fuzzing lautet der Fachbegriff für diese Art von Fuzz-Testing.

c't: Wenn das Target abstürzt, hat man einen Treffer gelandet?

Huster: Genau. Ein Absturz oder anderes unerwartetes Verhalten, zum Beispiel, wenn es sich aufhängt, sind eigentlich immer ein Hinweis auf einen Bug oder eine Schwachstelle. Die Eingaben, die so etwas bewirken, speichert der Fuzzer separat ab, sodass ich den Crash reproduzieren kann. Bei einer der fünf Lücken, die ich gefunden habe, war es zum Beispiel so, dass ein kaputtes Paket – oder eine Reihe von Paketen – eine sogenannte Linked List korrumpierte und quasi das letzte Paket in der Liste wieder auf das erste gezeigt hat. Bei der Verarbeitung wusste das Betriebssystem nie, wann die Liste zu Ende ist und hat sich schließlich aufgehängt, weil es aus dieser Schleife nicht rauskam.

c't: Das klingt nach einem ärgerlichen Bug, aber nicht nach einem, den ein Angreifer für eine Remote Code Execution nutzen könnte.

Huster: Nein. Es wäre schwierig, eine Möglichkeit zu finden, das auszunutzen. Die Endlosschleife führt dazu, dass das Betriebssystem sich aufhängt und das wars. Aber eine andere der Lücken ermöglicht es einem Angreifer, den Speicher zu überschreiben, sodass er theoretisch Code aus der Ferne

ausführen könnte. Der Kernel reserviert Speicher für die Ausführung von Programmen und Prozessen. Wenn jetzt beispielsweise 128 Byte an einer Stelle im Speicher für einen bestimmten Vorgang vorgesehen sind, dann darf man da eigentlich auch nicht mehr als diese 128 Byte reinschreiben. Bestimmte Eingaben des Fuzzers haben Fehler in der Paketverarbeitung aufgedeckt, die dazu führen, dass man mehr als die vorgesehene Länge in einen für einen Vorgang reservierten Teil des Speichers schreiben kann – ein sogenannter Buffer Overflow.

c't: Das wäre bereits ausreichend, damit ein Angreifer einen Rechner aus der Ferne übernehmen könnte?

Huster: Theoretisch. Es war möglich, als Angreifer 256 Byte kontrolliert in den Speicherbereich zu schreiben, der auf den zugewiesenen folgte. Für eine RCE müsste man zusätzlich herausfinden, wo im Speicher die kaputten WLAN-Pakete, die diesen Fehler im Kernel-Code triggern, überhaupt verarbeitet werden. Das ist aber gar nicht so einfach, weil es Mechanismen gibt, die dafür sorgen, dass der Kernel immer an unterschiedlichen Stellen im Speicher ausgeführt wird. Kernel Address Space Layout Randomization nennt sich das. Aber es wäre denkbar, dass sich noch weitere Sicherheitslücken finden, die einem das verraten.

c't: Ist das eine Hypothese oder hast du das auch erfolgreich prüfen können?

Huster: Nein. Das übersteigt meine Fähigkeiten. Es ist schon eher eine Hypothese. Aber eine, die sehr wahrscheinlich zutrifft. Es gibt verschiedene Arten von Sicherheitslücken und eine Lücke von diesem Typ bietet sich – in diesem konkreten Fall eben in Kombination mit weiteren – theoretisch dafür an.

Aus Angreifersicht das Spannende an den Sicherheitslücken ist, dass man überhaupt keine Nutzerinteraktion braucht. Du musst dich nicht aus Versehen mit einem Hotspot verbinden, den der

Hacker kontrolliert, damit er dir böse WLAN-Pakete schicken kann. Es reicht, wenn du dein WLAN anhast und dein Gerät nach Netzwerken in der Umgebung sucht. Im Hintergrund passiert das relativ häufig zur Standortbestimmung. Es ist nicht wie bei einem Phishing-Versuch, bei dem der Angreifer das Opfer erst dazu bringen muss, auf einen Button zu klicken und Login-Daten einzugeben. Genau das macht solche Lücken potenziell so kritisch. Linux-Nutzer gibt es nicht so viele, aber drei der Lücken betreffen Android, und Android-Nutzer gibt es eine ganze Menge. Am Smartphone haben die meisten Nutzer ihr WLAN in der Regel an.

c't: Ist der Fuzzer eine Eigenentwicklung des Secure Mobile Networking Labs?

Huster: Ja. Wir nutzen Komponenten aus LibAFL. Das ist eine Bibliothek, die ein sehr gutes Grundgerüst mitbringt, aber die Architektur unseres Fuzzers unterscheidet sich stark von der bestehender Fuzzer.

c't: Kannst du sicher sein, dass es außer den fünf Lücken nicht noch weitere gibt?

Huster: Ich denke, man kann auf jeden Fall sagen, dass WLAN unter Linux durch unsere Arbeit ein bisschen sicherer geworden ist. Wir waren an Stellen im Kernel, wo meines Wissens nach noch nicht so viel gefuzzt wurde. Momentan gucken wir uns noch weitere Teile an und bisher haben wir nichts weiter gefunden. Aber hundertprozentige Sicherheit, dass es nicht noch mehr Bugs und Sicherheitslücken gibt, wird man nie haben. Es kann immer unvorhergesehene Eingaben geben, die einen Bug oder eine Sicherheitslücke offenlegen. Ein Angreifer kann sie genauso gut finden wie wir. Genau deshalb ist Fuzz-Testing so wichtig.

c't: Seit Oktober gibt es Patches. Wie und an wen hast du die Sicherheitslücken gemeldet?

Huster: Es gibt gefühlt 1000 Anlaufstellen für Linux-Sicherheitssachen, zum Beispiel eine Mailing-Liste aller

Hersteller irgendwelcher Linux-Distributionen. Dort hätte ich das melden können. Parallel hätte ich dann noch die Kernel-Leute informieren müssen. Ich hab mich entschieden, den Prozess an einen Hersteller abzugeben und habe mich an SUSE gewandt. Die SUSE-Leute haben Johannes Berg von Intel ins Boot geholt. Er ist der Maintainer des WLAN-Stacks unter Linux. Für mich war es superspannend, mit ihm in so einem engen Austausch zu stehen, während er die Patches für die beiden Sicherheitslücken, die ich initial an SUSE gemeldet hatte, geschrieben hat.

Er hat mir die Patches dann geschickt und ich habe meinen Fuzzer darauf angesetzt. So sind wir auf die drei weiteren Sicherheitslücken – und insgesamt noch ein paar weitere kleinere Bugs – gestoßen. Das Ganze hat ein paar Wochen gedauert. Als alle Patches fertig waren, hat SUSE alle anderen Hersteller im Geheimen informiert und man hat einen Zeitpunkt festgelegt, zu dem man die Öffentlichkeit über die Lücken informiert. Die Hersteller hatten bis dahin über eine Woche Zeit, entsprechende Updates rauszubringen. Überrascht hat mich, dass manche Hersteller ihre Updates erst mehrere Tage nach der Bekanntgabe der Lücken verteilt haben.

c't: C gilt als relativ unsichere Programmiersprache. Künftig soll es möglich sein, Kernel-Komponenten stattdessen in Rust zu schreiben. Hätte das deine Sicherheitslücken verhindert?

Huster: Sehr wahrscheinlich wären diese Lücken nicht aufgetreten, hätte man die Module in Rust geschrieben. Gerade die Geschichte, dass man Speicher überschreiben kann. Der Rust-Compiler hätte verhindert, dass die Kernel-Entwickler diesen Fehler überhaupt einbauen. Aber es gibt natürlich auch Fehler, die durch keine Programmiersprache der Welt verhindert werden.

c't: Gibt es etwas, was du Admins und Anwendern raten würdest?

Huster: Sicherheitsupdates immer schnell einzuspielen. Wie

gesagt, bis alle größeren Distributionen die Updates verteilt haben, hat es nach Veröffentlichung noch ein paar Tage gedauert. Gerade bei Android dauert es oft länger. Es kann einfach sein, dass die betreffende Sicherheitslücke schon eine Weile öffentlich ist, bis man als Nutzer ein Sicherheitsupdate bekommt. Deshalb sollte man Updates möglichst sofort installieren. Auch wenn es nervt. Aber dann holt man sich in der Zwischenzeit halt mal einen Kaffee. (kst@ct.de)

Weitere Infos: ct.de/yvwk

Fake-Shops erkennen und Schäden vermeiden



Niemals ausgeliefert

Beim digitalen Einkauf betrügerischen Läden auf den Leim zu gehen, ist ärgerlich und teuer. Mit ein paar Grundregeln und hilfreichen Tools vermeidet man das. Wir stellen sie vor und erklären, was im Schadensfall noch möglich ist.

Beim digitalen Einkauf betrügerischen Läden auf den Leim zu gehen, ist ärgerlich und teuer. Mit ein paar Grundregeln und hilfreichen Tools vermeidet man das. Wir stellen sie vor und erklären, was im Schadensfall noch möglich ist.

Von Nick Akinci

Über vier Millionen Deutsche sind schon einmal auf einen Fake-Shop hereingefallen. Das schätzt das von der Bundesregierung geförderte Marktbeobachtungsinstitut „Marktwächter digitale Welt“. Besonders häufig bieten solche Shops nach Angaben des Instituts Sportartikel, Elektronik sowie Haushaltsartikel,

Bekleidung und Fahrräder, aber auch Brillen und Schmuck.

Wir zeigen, wie Sie Ihnen unbekannte Shops anhand verlässlicher Kriterien und mit hilfreichen Tools auf Seriosität prüfen, wie Sie Zahlungen absichern und was Sie tun können, falls Sie doch auf einen Fake-Shop hereingefallen sind.

Was ist ein Fake-Shop?

Fake-Shops sind Online-Shops, mit denen Kriminelle gutgläubigen Kunden ihr Geld abnehmen wollen, ohne ihnen die versprochene Ware zu liefern. In der einfachsten Variante erhalten Kunden, die darauf hereinfallen, überhaupt keine Ware. Etwas perfidere Betrüger versenden leere Kartons. Im Nachhinein behaupten sie, dass die Ware auf dem Versandweg abhandengekommen sein müsse. Mitunter verschicken sie auch Ware, die in keiner Weise der Produktbeschreibung entspricht.

Viele Fake-Shops sind nur für einen relativ kurzen Zeitraum online, da sie fast immer auffliegen und der Hoster sie im besten Fall vom Netz nimmt. In diesem Zeitfenster versuchen die Betrüger, möglichst viel Geld zu ergaunern. Sitzt der Hoster im Ausland, können sich solche Shops auch über Jahre halten.

Prüfender Blick

Fake-Shops sind häufig nicht auf den ersten Blick als solche zu erkennen. In Zeiten von Baukastensystemen wie Shopify & Co. klicken Betrüger professionell aussehende Online-Shops in wenigen Stunden zusammen. Es gibt jedoch eine Reihe von Indizien, die für einen Fake-Shop sprechen.

Um Kunden anzulocken, bieten die Täter die Ware in Fake-Shops oft deutlich günstiger an als in anderen Online-Shops. Insbesondere beliebte und häufig gehandelte Markenware preisen sie unter dem Marktwert an, gern als Sonderangebot getarnt.

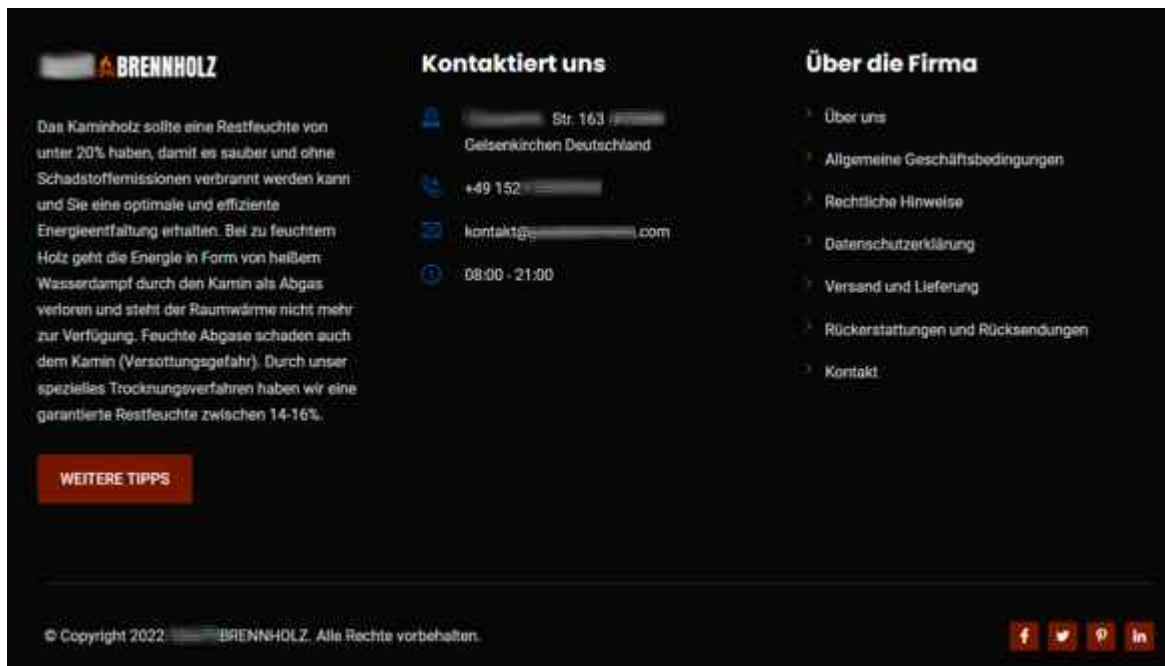
Schnäppchenjäger können sich auf Preisvergleichsseiten einen Eindruck verschaffen, ob die Preisgestaltung realistisch ist.

Als Nächstes schaut man in das Impressum. Fake-Shops haben oft keines, obwohl dies in Deutschland gesetzliche Pflicht ist – die Betrüger wollen ihre Identität verschleiern. Aber Achtung: Manche Fake-Shops enthalten ein echt aussehendes Impressum, welches jedoch schlicht falsche, unvollständige oder von anderen Websites kopierte Angaben enthält. Ob die Firma an der angegebenen Adresse sitzt, kontrolliert man am besten mit Google Maps. Den Unternehmensnamen und die zugehörige Handelsregisternummer prüft man auf handelsregister.de [1].

Abgesehen vom Impressum fehlen in vielen Fake-Shops auch Telefonnummern oder E-Mail-Adressen, um Kontakt aufzunehmen. Ebenfalls kein gutes Zeichen ist es, wenn sich Kontaktmöglichkeiten beschränken auf ausschließlich Handy- oder kostenpflichtige Nummern, Postfachadressen oder lediglich ein Kontaktformular. Misstrauen ist geboten, wenn AGB und Datenschutzerklärung sowie Widerrufsbelehrungen und Versandbedingungen fehlen.

Gütesiegel sind ein Hinweis auf vertrauenswürdige Shops, doch in Fake-Shops trifft man immer wieder einfach hineinkopierte oder frei erfundene Varianten an. Letztere ähneln teils bekannten Gütesiegeln – wie etwa dem von [Trusted Shops](https://www.trustedshops.de).

Verfügt der Online-Shop über ein Gütesiegel, kann man auf der Homepage der Organisation prüfen, ob es sich um ein tatsächlich anerkanntes Gütesiegel handelt und ob der Online-Shop es rechtmäßig erworben hat. Durch einen Klick auf das Siegelsymbol muss man auf die Seite der dahinterstehenden Organisation gelangen. Verbreitet und vertrauenswürdig ist außer Trusted Shops auch das [EHI Retail Institute](https://www.ehi-retail-institute.de) („Geprüfter Online-Shop“). Als zuverlässig gilt außerdem das in Kopenhagen ansässige Bewertungsportal [Trustpilot](https://www.trustpilot.com) (alle unter [ct.de/yu3d](https://www.ct.de/yu3d)).



Kein Impressum, kein Handelsregistereintrag, keine Umsatzsteuer-ID, Shop ganz neu, Google Maps kennt den Shop an der angegebenen Adresse nicht und als Kontaktmöglichkeit nur eine Mobiltelefonnummer: Hier heißt es Finger weg!

Zahlungsmethoden

Als Zahlart bieten viele Fake-Shops ausschließlich Vorkasse per Banküberweisung an, da man solche Zahlungen in der Regel nicht rückgängig machen kann. Mitunter wollen betrügerische Händler Kunden auch gerne zu PayPal-Zahlungen in der Variante „Freunde und Familie“ verleiten. Die beinhalten aber im Unterschied zur Option „Waren und Dienstleistungen“ keinen Käuferschutz. Manchmal bietet der Fake-Shop auch zum Schein weitere Zahlarten an, um Vertrauen zu schaffen. Die funktionieren dann aber aus vorgeschobenen Gründen nicht. Daraufhin bitten die Täter um Vorkasse oder die unsichere PayPal-Variante.

Auch bei vermeintlich sicheren Bezahlungsmethoden gibt es Haken. Der PayPal-Käuferschutz ist zum Beispiel an Bedingungen wie Paketversand mit elektronischer Sendungsverfolgung geknüpft [3]. Ähnlich halten es Amazon oder Klarna. Manche Betreiber von Fake-Shops schicken die Pakete daher an Adressen von Strohleuten, um Kunden über die Sendungsverfolgung erst in

Sicherheit zu wiegen und anschließend Käuferschutzverfahren zu erschweren. Mehr zu Vor- und Nachteilen von Zahlarten haben wir unter [2] zusammengetragen.

Blacklists und Prüftools

Bleibt man unsicher, helfen Tools von Verbraucherschützern und anderen Organisationen. Zunächst lohnt sich ein Blick auf Blacklists. Hierbei handelt es sich um Listen von Online-Shops, die bereits als Fake-Shops eingestuft oder die mehrfach als solche gemeldet worden sind. Solche Listen finden sich zum Beispiel auf der [Website der Verbraucherzentrale Hamburg](#), der [Präsenz des Siegel-Anbieters Trusted Shops](#) oder auf der [Watchlist Internet](#). Der [Fake-Shop-Kalender](#) der Verbraucherzentrale Bundesverband macht zusätzlich auf zeitweise besonders häufig betroffene Branchen aufmerksam (alle Seiten unter ct.de/yu3d). Darüber hinaus kann sich der Besuch der Preisvergleichsseiten Geizhals und Idealo lohnen (Hinweis: Geizhals gehört wie c't zu Heise Medien). Sie listen nur geprüfte Online-Shops sowie Händler auf Marktplätzen mit starkem Käuferschutz. Mehr zu den Eigenheiten von Marktplätzen wie Amazon und eBay finden Sie unter [3].



Fakeshop-Finder

Ist dieser Online-Shop seriös?

kramerversand.de	Shop-URL prüfen
------------------	-----------------

Diese Shop-URL weist Anzeichen für einen Fakeshop auf.



Einschätzung:

Zu diesem Shop liegen mehrere Anzeichen für einen Fakeshops vor. Der Fakeshop-Finder konnte das Impressum des Shops nicht auslesen. Das kann beispielsweise passieren, wenn die entsprechenden Seiten von den Shops für automatisierte Anfragen gesperrt wurden. Das heißt nicht, dass es sich um einen Fakeshop handelt. Bitte [überprüfen Sie in diesem Fall selbst](#), ob Sie ein Impressum auf den Seiten finden können.

Wichtige Fakeshop-Merkmale:

- ✗ Es wurde kein Impressum gefunden.
Der Fakeshop-Finder konnte automatisch kein Impressum finden. Das kann beispielsweise passieren, wenn die entsprechenden Seiten von den Shops für automatisierte Anfragen gesperrt wurden. Bitte überprüfen Sie in diesem Fall selbst, ob Sie ein Impressum auf den Seiten - meistens im unteren Bereich - finden können.
- ✗ Fakeshop Warnungen:
 - Dieser Online-Shop wurde am 20.08.2022 von seitcheck.de als Fakeshop eingestuft. Zum Eintrag bei [seitcheck.de](#)
 - Dieser Online-Shop wurde am 19.08.2022 von auktionenhilfe.info als Fakeshop eingestuft. Zum Eintrag bei [auktionshilfe.info](#)
 - Dieser Online-Shop wurde am 22.08.2022 von Watchlist Internet als Fakeshop eingestuft. Zum Eintrag bei [Watchlist Internet](#)
 - Dieser Online-Shop wurde am 22.08.2022 von Trusted Shops als Fakeshop eingestuft. Zum Eintrag bei [Trusted Shops](#)

Mit dem Fakeshop-Finder der Verbraucherzentralen überprüft man Shop-Websites. Bei einer roten Ampel handelt es sich nahezu sicher um einen Fake-Shop.

Hilfreich bei der Recherche ist außerdem der [Fakeshop-Finder](#) der Verbraucherzentralen. Dort gibt man die URL des zu prüfenden Online-Shops in eine Eingabemaske ein. Anschließend ordnet das Tool ihn nach einem Ampelsystem einer Kategorie zu. Zeigt die Ampel Rot, so ist der betreffende Shop bereits als Fake-Shop aufgefallen. Bei gelber Ampelfarbe hat die automatische Prüfung allgemeine Indizien für betrügerische Absichten, aber auch Indizien für seriöses Gebaren gefunden und listet sie samt Erklärung auf. Entdeckt die Prüfroutine beispielsweise kein Impressum, kann das auch heißen, dass der Betreiber des Shops es lediglich für automatisierte Abfragen gesperrt hat. Das muss man dann selbst nachsehen. Die Einstufung „Grün“ bedeutet, dass der Shop den

Verbraucherzentralen „bisher nicht negativ aufgefallen“ ist; man soll aber trotzdem auf eine sichere Zahlungsmethode und die Rücksendekonditionen achten.

Schäden begrenzen, Shops melden

Ist das Kind bereits in den Brunnen gefallen, kann man versuchen, das im Fake-Shop ausgegebene Geld zurückzubekommen. Im besten Fall hat man eine sichere Zahlungsmethode verwendet und veranlasst über seine Bank oder den Zahlungsdienstleister eine Rückerstattung. Bei einer Banküberweisung wird es hingegen schwierig. Meldet man sich sofort oder zumindest am selben Tag bei seiner Bank, kann diese die Überweisung manchmal noch stoppen.

In jedem Fall sollte man Strafanzeige bei der Polizei oder Staatsanwaltschaft erstatten. Dies geht heutzutage unkompliziert über die „[Onlinewache](https://www.ct.de/yu3d)“ ([ct.de/yu3d](https://www.ct.de/yu3d)). Zusätzlich kann man einen Rechtsanwalt damit beauftragen, den Rückzahlungsanspruch auf zivilrechtlicher Ebene durchzusetzen. Der Anwalt beantragt Einsicht in die Ermittlungsakte der Strafverfolgungsbehörden und findet im besten Fall die Identität des Betrügers heraus.

Wer einen Fake-Shop erkannt hat oder darauf hereingefallen ist, kann dazu beitragen, dass der Shop aus dem Internet verschwindet. Hat man als Betroffener Strafanzeige erstattet, kümmern sich meist Polizei und Staatsanwaltschaft darum, dass der Hoster den Shop abschaltet. Ansonsten meldet man den Fake-Shop dem Hoster oder Shopsystemanbieter sowie den Verbraucherzentralen, zum Beispiel über das [Onlineformular der Verbraucherzentrale Hamburg](https://www.ct.de/yu3d) ([ct.de/yu3d](https://www.ct.de/yu3d)). (mon@ct.de)

1. Literatur
2. [Jo Bager, Gefährliche Offenheit, Online-Handelsregister lädt zum Datenmissbrauch ein, c't 24/2022, S. 134](#)
3. [Markus Montz, Geld her!, Onlinekauf-Checkliste](#)

[Bezahlmethoden, c't 8/2022, S. 26](#)

4. [Georg Schnurer, Händler-Roulette, Onlinekauf-Checkliste Shop-Auswahl, c't 8/2022, S. 24](#)

Nützliche Websites: ct.de/yu3d

SQL Injection in Java mit JPA und Hibernate verhindern



entwickler.de – entwickler.de Deine Wissensplattform

[...]Weiterlesen...

Wirft man einen Blick auf die Top-10-Schwachstellen der OWASP [1], sind SQL Injections immer noch in einer prominenten Position zu finden. In diesem Artikel diskutieren wir verschiedene Möglichkeiten, wie SQL Injections effizient vermieden werden können.

Wenn Anwendungen auf Datenbanken zugreifen, bestehen immer wieder hohe Sicherheitsrisiken für die Applikation. Hat ein Angreifer die Möglichkeit, die Datenbankschicht einer Anwendung zu kapern, kann er zwischen mehreren Optionen wählen. Die Daten der gespeicherten Benutzer zu stehlen, um sie mit Spam zu überfluten, ist dabei nicht das schlimmste mögliche Szenario. Noch problematischer wäre es, wenn gespeicherte Zahlungsinformationen missbraucht würden. Eine weitere Variante eines SQL-Injection-Cyberangriffs ist der illegale Zugriff auf eingeschränkte kostenpflichtige Inhalte

und/oder Dienste. Wie wir sehen, gibt es viele Gründe, sich um die Sicherheit von (Web-)Anwendungen zu kümmern.

Um eine gut funktionierende Prävention gegen SQL Injections etablieren zu können, müssen wir zunächst verstehen, wie ein solcher Angriff funktioniert und auf welche Punkte wir achten müssen. Kurz gesagt verhält es sich so: Jede Benutzerinteraktion, die die Eingabe ungefiltert in einer SQL-Abfrage verarbeitet, ist ein mögliches Angriffsziel. Die Dateneingabe kann so manipuliert werden, dass die übermittelte SQL-Abfrage eine andere Logik enthält als das Original. Der folgende Code gibt eine gute Vorstellung davon, was möglich ist:

```
SELECT Username, Password, Role FROM User
  WHERE Username = 'John Doe' AND Password = 'S3cr3t';
SELECT Username, Password, Role FROM Users
  WHERE Username = 'John Doe'; --' AND Password='S3cr3t';
```

Die erste Anweisung zeigt die ursprüngliche Abfrage. Wird die Eingabe für die Variablen Benutzername und Passwort nicht gefiltert, entsteht das klassische Angriffsszenario. Die zweite Abfrage fügt für die Variable Benutzername einen String mit dem Benutzernamen *John Doe* ein und erweitert ihn um die Zeichen ; –. Diese Anweisung umgeht die *UND*-Verzweigung und gibt in diesem Fall Zugriff auf das Log-in. Mit der Zeichensequenz ,; schließen Sie die *WHERE*-Anweisung und mit – werden alle folgenden Zeichen auskommentiert. Theoretisch ist es möglich, zwischen diesen beiden Zeichenfolgen jeden gültigen SQL-Code auszuführen. Es lässt sich leicht ahnen, welcher Schabernack an dieser Stelle möglich ist.

Mein Plan ist natürlich nicht, zu verbreiten, welche SQL-Befehle die schlimmsten Folgen für das Opfer haben könnten. Bei diesem einfachen Beispiel gehe ich davon aus, dass die Botschaft klar angekommen ist. Wir müssen jede UI-Eingabevariable in unserer Anwendung vor Benutzermanipulation schützen. Auch dann, wenn sie nicht direkt für Datenbankabfragen verwendet werden. Um diese Variablen zu

erkennen, ist es immer eine gute Idee, alle vorhandenen Eingabeformulare zu validieren. Doch moderne Anwendungen haben meist mehr als nur ein paar Eingabeformulare. Aus diesem Grund sage ich auch sehr eindringlich: Behalten Sie Ihre REST-Endpunkte im Auge. Oft sind deren Parameter auch mit SQL-Abfragen verbunden.



Security Afternoon

Durch einen stetigen Strom an Releases, neuen Features und spannenden Projekten rückt Security in der IT-Welt gerne einmal in den Hintergrund. Beim Security Afternoon rücken wir mit Michael Kaufmann und Inko Lorch einen ganzen Nachmittag lang die IT-Sicherheit in den Fokus und zeigen, warum es so wichtig ist, Anwendungssicherheit nicht als lästige Fleißaufgabe zu verstehen.

Deshalb sollte die Eingabevalidierung generell Teil des Sicherheitskonzepts sein. Annotationen aus der Spezifikation Bean Validation [2] sind für diesen Zweck sehr mächtig. Beispielsweise sorgt `@NotNull` als Annotation für das Datenfeld im Domänenobjekt dafür, dass das Objekt nur persistiert werden kann, wenn die Variable nicht leer ist. Um die Bean Validation Annotations in Ihrem Java-Projekt zu verwenden, müssen Sie

lediglich eine kleine Bibliothek einbinden:

```
<dependency>
  <groupId>org.hibernate.validator</groupId>
  <artifactId>hibernate-validator</artifactId>
  <version>${version}</version>
</dependency>
```

Eventuell ist es notwendig, komplexere Datenstrukturen zu validieren. Mit regulären Ausdrücken haben Sie ein weiteres mächtiges Werkzeug an der Hand. Aber seien Sie vorsichtig: Es ist nicht so einfach, korrekt funktionierende RegEx zu schreiben. Schauen wir uns dazu ein kurzes Beispiel an (Listing 1).

Listing 1: Validierung durch reguläre Ausdrücke in Java

```
public static final String RGB_COLOR = "#[0-9a-fA-F]{3,3}([0-9a-fA-F]{3,3})?";
```

```
public boolean validate(String content, String regEx) {
    boolean test;
    if (content.matches(regEx)) {
        test = true;
    } else {
        test = false;
    }
    return test;
}
```

```
validate('#000', RGB_COLOR);
```

Die RegEx zur Erkennung des korrekten RGB-Farbschemas ist recht einfach. Gültige Eingaben sind `#fff` oder `#000000`. Der Bereich umfasst die Zeichen `0-9` und zusätzlich noch Buchstaben `A-F`. Groß-/Kleinschreibung wird in unserem Beispiel nicht beachtet. Wenn Sie Ihre eigene RegEx entwickeln, müssen Sie bestehende Grenzen immer sehr gut im Auge behalten. Ein gutes Beispiel, um obere beziehungsweise untere Schranken zu verstehen, ist das 24-Stunden-Zeitformat. Typische Fehler sind ungültige Eingaben wie `23:60` oder `24:00`. Ein Blick auf die

Anzeige der Digitaluhr zeigt für ein 24-Stunden-Format als untere Schranke `00:00` und als obere Schranke `23:59`, alles andere ist ungültig.

Die Methode `validate` vergleicht die Eingabezeichenfolge mit der RegEx. Wenn das Muster mit der Eingabe übereinstimmt, gibt die Methode `TRUE` zurück. Wenn Sie weitere Ideen zu Validatoren in Java erhalten möchten, können Sie auch in meinem GitHub-Repository [3] nachsehen.

Zusammengefasst ist unsere erste Idee, um Benutzereingaben vor Missbrauch zu schützen, alle problematischen Zeichenfolgen herauszufiltern wie SQL-Kommentare und so weiter. Und solch eine Sperrliste ist auch nicht schlecht. Zumindest für den Anfang. Eine Blacklist weist aber einige Einschränkungen auf. Zunächst erhöht sich die Komplexität der Anwendung, da das Blockieren einzelner Zeichen wie `-;` und `,` manchmal unerwünschte Nebenwirkungen verursachen kann. Auch eine anwendungsweite Standardbegrenzung der Zeichen könnte Probleme bereiten. Stellen Sie sich vor, es gibt einen Textbereich für ein Blogsystem oder Ähnliches.

Das bedeutet, dass wir ein weiteres leistungsstarkes Konzept benötigen, um die Eingabe so zu filtern, dass unsere SQL-Abfrage nicht manipuliert werden kann. Um dieses Ziel zu erreichen, bietet der SQL-Standard eine sehr gute Lösung. SQL-Parameter sind Variablen innerhalb einer SQL-Abfrage, die als Inhalt und nicht als Anweisung interpretiert werden. Das ermöglicht es, große Texte entgegenzunehmen, ohne einige gefährliche Zeichen blockieren zu müssen. Schauen wir uns an, wie das mit einer PostgreSQL-Datenbank [4] funktioniert:

```
DECLARE user String;  
SELECT * FROM login WHERE name = user;
```

Für den Fall, dass Sie den OR-Mapper Hibernate [5] verwenden, gibt es mit dem Java Persistence API (JPA) einen eleganteren Weg (Listing 2).

Listing 2: Hibernate-JPA-SQL-Parameter verwenden

```
String myUserInput;

@PersistenceContext
public EntityManager mainEntityManagerFactory;

CriteriaBuilder builder =
    mainEntityManagerFactory.getCriteriaBuilder();

CriteriaQuery<DomainObject> query =
    builder.createQuery(DomainObject.class);

// create Criteria
Root<ConfigurationD0> root =
    query.from(DomainObject.class);

//Criteria SQL Parameters
ParameterExpression<String> paramKey =
    builder.parameter(String.class);

query.where(builder.equal(root.get("name"), paramKey));

// wire queries together with parameters
TypedQuery<ConfigurationD0> result =
    mainEntityManagerFactory.createQuery(query);

result.setParameter(paramKey, myUserInput);
DomainObject entry = result.getSingleResult();
```

Listing 2 zeigt ein vollständiges Beispiel für Hibernate mit JPA und dem Criteria API. In der ersten Zeile wird die Variable für die Benutzereingabe deklariert. Die Kommentare in der Auflistung erklären sehr deutlich, wie es funktioniert. Wie Sie sehen können, ist das keine Raketenwissenschaft. Die Lösung hat neben der Verbesserung der Sicherheit von Webanwendungen einige weitere nette Vorteile. So wird kein einfaches SQL verwendet. Dadurch wird sichergestellt, dass jedes Datenbankverwaltungssystem, das von Hibernate unterstützt wird, durch diesen Code gesichert werden kann.

Die Nutzung sieht vielleicht etwas komplizierter aus als eine einfache Abfrage, aber der gewonnene Nutzen für Ihre Anwendung ist enorm. Andererseits gibt es natürlich einige zusätzliche Codezeilen. Aber die sind nicht so schwer zu verstehen, wie dieser Artikel gezeigt hat.



Marco Schulz studierte an der HS Merseburg Diplominformatik und twittert regelmäßig als @ElmarDott über alle möglichen technischen Themen. Seine Schwerpunkte sind hauptsächlich Build- und Konfigurationsmanagement, Softwarearchitekturen und Release-Management. Seit knapp 20 Jahren realisiert er in internationalen Projekten für namhafte Unternehmen umfangreiche Webapplikationen. Er ist freier Consultant/Trainer. Sein Wissen teilt er mit anderen Technikbegeisterten auf Konferenzen, wenn er nicht gerade wieder einmal an einem neuen Fachbeitrag schreibt.

Links & Literatur

[1] <https://owasp.org>

[2] <https://beanvalidation.org>

[3]

<https://github.com/ElmarDott/TP-CORE/blob/master/src/main/java/org/europa/together/utils/Validator.java>

[4]

<https://www.postgresql.org/docs/9.1/plpgsql-declarations.html>

[5] <https://hibernate.org>

[6] <https://elmar-dott.com/courses/de/web-application-security>

[7]

Originalartikel:

<https://elmar-dott.com/articles/preventing-sql-injections-in-java/>

TLS mit Wireshark entschlüsseln



TLS mit Wireshark entschlüsseln

Was es beim kriminellen Man in the Middle zu verhindern gilt, gehört bei legal agierenden Systemadmins zum notwendigen Handwerkszeug: der Zugriff auf verschlüsselte Datenströme zwecks Fehlersuche.

Was es beim kriminellen Man in the Middle zu verhindern gilt, gehört bei legal agierenden Systemadmins zum notwendigen Handwerkszeug: der Zugriff auf verschlüsselte Datenströme zwecks Fehlersuche.

Von Benjamin Pfister

Der Anteil des verschlüsselten Datenverkehrs nimmt ständig zu. Fast alle Webdienste nutzen Transport Layer Security (TLS) und aktuelle Browser warnen bei unverschlüsselten HTTP-

Verbindungen ausdrücklich vor dem damit verbundenen Risiko. Das ist aus Sicht der Sicherheit und des Datenschutzes sehr zu begrüßen – doch die Verschlüsselung verhindert auch eine legale Analyse des Datenstroms, etwa seitens berechtigter Admins. Es gibt jedoch Möglichkeiten der Fehlersuche trotz TLS-Verschlüsselung, zum Beispiel mit dem im Folgenden beschriebenen Paketanalysewerkzeug Wireshark.

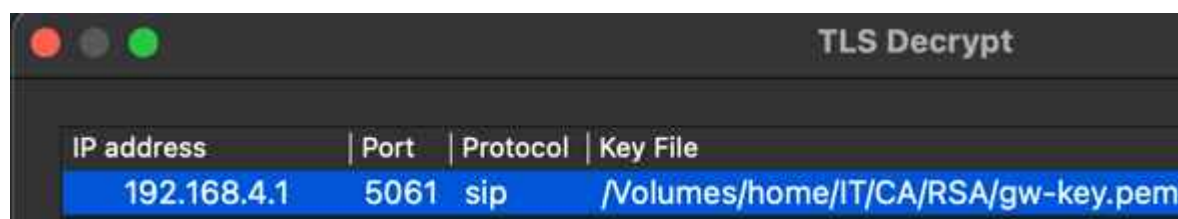
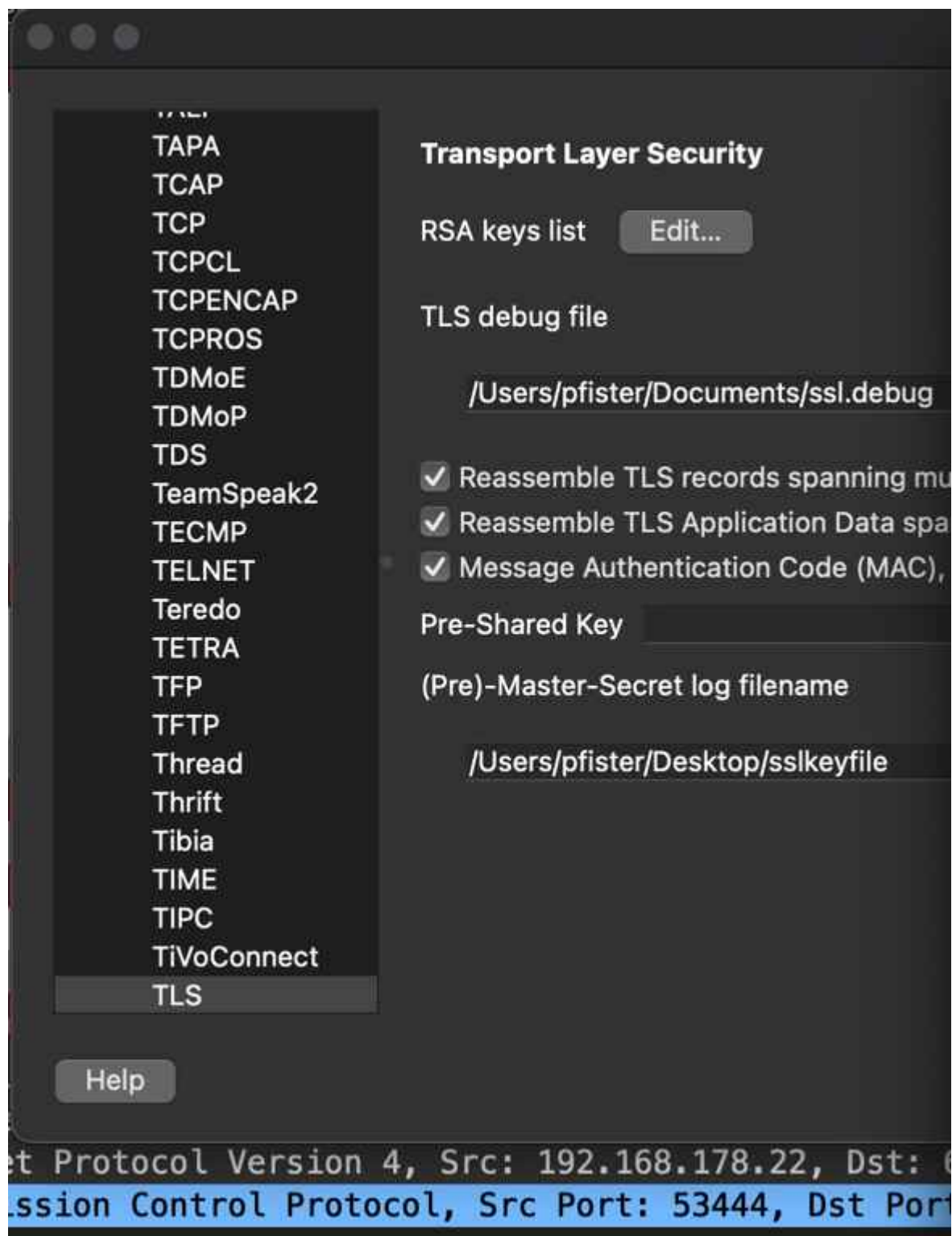
Wireshark bringt einen eigenen Dissector (wörtlich übersetzt Sezierer) für TLS mit. Er ermöglicht neben der Aufteilung und Darstellung der Protokolle auch die Entschlüsselung der Nutzdaten. Dazu bedarf es der passenden Schlüssel. Je nach eingesetzter Cipher Suite kommen unterschiedliche Entschlüsselungsmethoden zum Einsatz: auf Basis eines Session- (Pre-Master Secret) oder eines privaten RSA-Schlüssels.

Welche der beiden zur Anwendung kommt, hängt von der Cipher Suite ab: mit Perfect Forward Secrecy (PFS) oder ohne. Falls für die Übertragung keine PFS Cipher Suites vorgesehen sind, kann die Entschlüsselung auf Basis des privaten Schlüssels des Serverauthentifizierungszertifikats stattfinden. In diesem Fall kann Wireshark jedoch auch die Methode Pre-Master Secret nutzen. Dies ist beispielsweise dann interessant, wenn man – wie bei öffentlichen Webdiensten – nicht im Besitz der privaten Schlüssel ist.

Bei Nutzung von Perfect Forward Secrecy (PFS) lässt sich der Datenstrom selbst bei Kenntnis des privaten Schlüssels nicht nachträglich entschlüsseln. Daher empfiehlt das BSI zum Schutz personenbezogener oder anderer sensibler Daten diese Cipher Suites. Darunter fallen die Cipher Suites mit Diffie-Hellman Ephemeral (DHE) und Elliptic Curve Diffie-Hellman Ephemeral (ECDHE). Um diese Varianten zu entschlüsseln, muss man die Methode Pre-Master Secret einsetzen.

Der Besitz des privaten Schlüssels nützt also nur dann etwas, wenn keine (EC)DHE Cipher Suites zum Einsatz kommen. Zudem funktioniert diese erste Variante nicht mit TLS 1.3. Einen

weiteren Fallstrick birgt der TLS Session Resume, bei dessen Anwendung das Entschlüsseln fehlschlägt. Es bedarf der Aufzeichnung eines ClientKeyExchange im TLS Handshake. Zum Entschlüsseln der Daten benötigt man das Serverauthentifizierungszertifikat – genauer dessen privaten Schlüssel. In den TLS-Protokolleinstellungen von Wireshark und dem Menüpunkt „RSA keys list“ referenziert man die Datei mit dem privaten Schlüssel und verknüpft ihn mit der IP-Adresse, dem Port und dem Protokoll des Servers. Abbildung 1 zeigt eine solche Hinterlegung für die IP-Adresse 192.168.4.1 mit dem Port 5061 und dem Protokoll SIP. Der Private Key liegt im Beispiel unter /Volumes/Home/IT/CA/RSA/gw-key.pem. Daran erkennt man, dass nicht nur HTTPS als Applikationsprotokoll zur Verfügung steht. Die Referenzen liegen im Beispiel von macOS unter /Users/<username>/.config/wireshark/ssl_keys.



Hinterlegung des Private Key aus dem Serverauthentifizierungszertifikat (Abb. 1).

Nach der korrekten Hinterlegung beginnt der Dissector mit der Entschlüsselung. Bei eventuellen Fehlern lohnt ein Blick in

die TLS-Debug-Datei, die beispielsweise fehlerhafte Private-Key-Zuweisungen oder Probleme beim Laden der Private Keys aufzeigt. Deren Zielverzeichnis und Namen kann man selbst wählen (siehe Abbildung 1).

Auf der Kommandozeile kann man das in Wireshark enthaltene CLI-Tool tshark nutzen. Für die RSA-Methode lautet der Befehl

```
tshark -o "ssl.keys_list:
192.168.4.1,5061,sip,/Volumes/Home/IT/CA/RSA/gw-key.pem" -r
siptls.pcapng -Y sip
```

Über die Option

```
-o "ssl.keys_list:
192.168.4.1,5061,sip,/Volumes/Home/IT/CA/RSA/gw-key.pem"
```

verknüpft man die in Abbildung 1 dargestellten Einstellungen – ähnlich wie mit der GUI-Variante. Das Argument -r siptls.pcapng liest dabei lediglich die PCAPNG-Datei. Das Argument -Y sip setzt einen Display-Filter auf das VoIP-Signalisierungsprotokoll SIP, sodass keine Pakete anderer Protokolle die Ausgabe fluten.

Die zweite Variante – keine Kenntnis des privaten Schlüssels und der Einsatz von (EC)DHE – setzt eine Keylog-Datei voraus, also eine Textdatei, die von unterschiedlichen Kryptobibliotheken bereitgestellt wird, beispielsweise OpenSSL oder NSS. Darauf aufbauende Applikationen wie Chrome, Firefox oder Curl generieren diese Datei, wenn die Umgebungsvariable SSLKEYLOGFILE gesetzt ist. Unter macOS kann man diese beispielsweise wie folgt anlegen: export SSLKEYLOGFILE="/Users/<username>/Desktop/sslkeyfile".

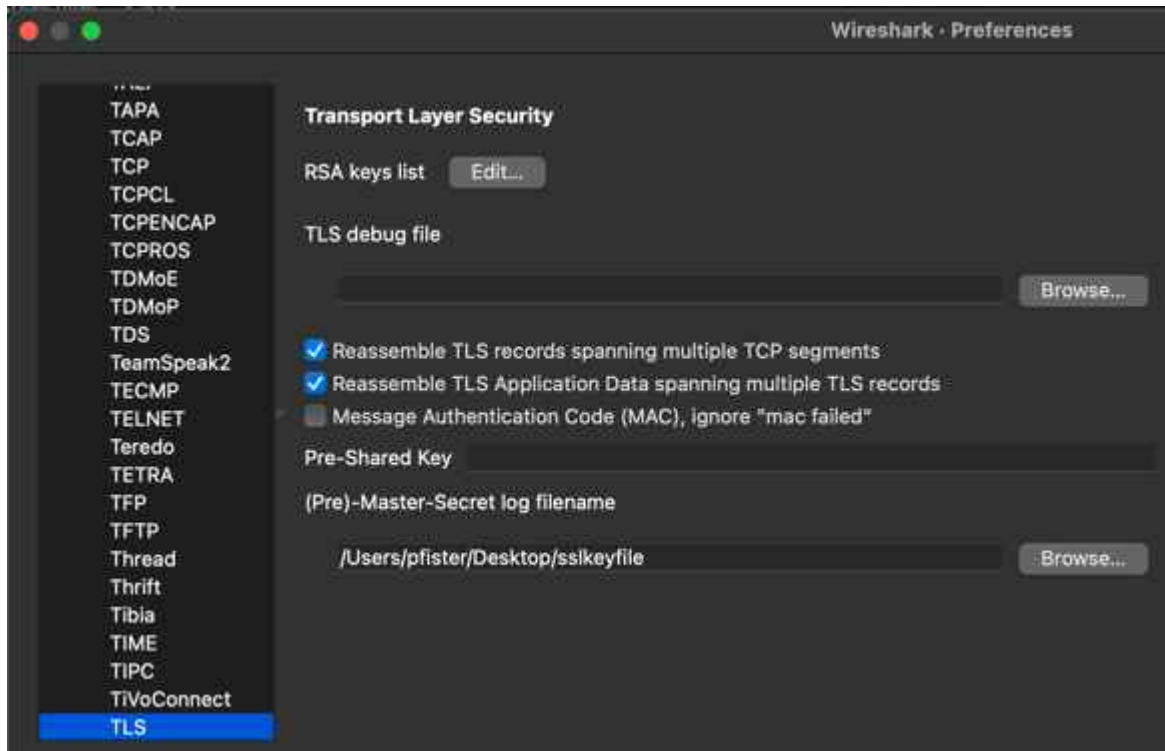
Die Bibliotheken schreiben den Pre-Master Key dann in die in der Umgebungsvariablen referenzierte Datei. Der Client generiert diesen in der Client Exchange Phase des TLS Handshake. Der Export kann auf Client- oder Serverseite stattfinden. Ein Mitlesen auf dem Transportweg ist somit nicht möglich. Wireshark kann den Pre-Master Key aus dem Handshake

dafür nutzen, den Master Key abzuleiten und damit den Datenverkehr zu entschlüsseln. Im Anschluss an die Konfiguration der Umgebungsvariablen startet man den Mitschnitt in Wireshark und öffnet dann über die Konsole beispielsweise Firefox mittels `open /Applications/Firefox.app` unter macOS. Nachdem die erste TLS-verschlüsselte Seite aufgerufen wurde, zeigt sich, ob die Schlüsseldatei korrekt gespeichert wurde. In der ersten Zeile der Datei erkennt man auch, dass die Bibliothek NSS für den Schreibvorgang zuständig war (siehe Abbildung 2).

```
pfister@wlc:~$ cat Desktop/sslkeyfile
# SSL/TLS secrets log file, generated by NSS
CLIENT_HANDSHAKE_TRAFFIC_SECRET d9f54f9b831c8a29ee6438f4a86e6277f84663ba25f32bd08e466ce82d18488 75bcb0fef64e3386f7d2a23c39e98c45a77f84666b627138b1d880fca7e5V
a4e
SERVER_HANDSHAKE_TRAFFIC_SECRET d9f54f9b831c8a29ee6438f4a86e6277f84663ba25f32bd08e466ce82d18488 19d7275d9ee9fff7c615228e38f73a8ac29f93c235d67a34aa3788183c89e6
13a
CLIENT_TRAFFIC_SECRET_0 d9f54f9b831c8a29ee6438f4a86e6277f84663ba25f32bd08e466ce82d18488 97c8432787a3d944c13e8a338d137adfe761fde21885e81a9c869304a461619
SERVER_TRAFFIC_SECRET_0 d9f54f9b831c8a29ee6438f4a86e6277f84663ba25f32bd08e466ce82d18488 88968481e54e71a2e3a3748b6732c2a48e8085199b384a8a973e7284392d6b0
EXPORTER_SECRET d9f54f9b831c8a29ee6438f4a86e6277f84663ba25f32bd08e466ce82d18488 b29eb720a38e3a18546c6eccfa2251b4e2cb3e18c46b52288b6df1272fe0fc
CLIENT_HANDSHAKE_TRAFFIC_SECRET 548a1296b77b22a686c5970184b1391a6ef7b5f2be5a6e3fa368838589a9c5 c378c843a2d2a8d81eadf8c3638e462942d535a12d046f86722823dc7456
4ed
SERVER_HANDSHAKE_TRAFFIC_SECRET 548a1296b77b22a686c5970184b1391a6ef7b5f2be5a6e3fa368838589a9c5 c8216553efbc3780ec52b1956f37ef5c82849764a9e95d67a4d8689e6c63
a37
CLIENT_TRAFFIC_SECRET_0 548a1296b77b22a686c5970184b1391a6ef7b5f2be5a6e3fa368838589a9c5 815c38ea95e98673b78205eb4323e7344896eb2ef86cd9890388551624868
SERVER_TRAFFIC_SECRET_0 548a1296b77b22a686c5970184b1391a6ef7b5f2be5a6e3fa368838589a9c5 9e776a8ba98a3bc38a87c558c1f1fa11de7c977a418e1865d34579aebdfca73a
EXPORTER_SECRET 548a1296b77b22a686c5970184b1391a6ef7b5f2be5a6e3fa368838589a9c5 24d1a864a980dfc77b82ef8b7648and3e2115bde7e3ba1d6c98b89f8e871c062b
CLIENT_HANDSHAKE_TRAFFIC_SECRET 997a334266ba1a91eb07ae4e69ea72a7842f77e672a8d7ea33bac783314744 f91ba3841d91ce886f2198013b6739cf8fe75fa3a9f7d3caa673119e0b8
f4e
SERVER_HANDSHAKE_TRAFFIC_SECRET 997a334266ba1a91eb07ae4e69ea72a7842f77e672a8d7ea33bac783314744 f91ba3841d91ce886f2198013b6739cf8fe75fa3a9f7d3caa673119e0b8
f4e
CLIENT_HANDSHAKE_TRAFFIC_SECRET 51f19abff428V3a59c0fe7ebcac84b3da2589594a4687d386796c6446356799 bf5856e31a8aaa498a79ba853fc8cd52760bcb97c48c4299791a853cdca4
f07
SERVER_HANDSHAKE_TRAFFIC_SECRET 51f19abff428V3a59c0fe7ebcac84b3da2589594a4687d386796c6446356799 bf5856e31a8aaa498a79ba853fc8cd52760bcb97c48c4299791a853cdca4
82e
CLIENT_TRAFFIC_SECRET_0 997a334266ba1a91eb07ae4e69ea72a7842f77e672a8d7ea33bac783314744 bffba13231a2d8eb5a5ff9d6078285de1dc218b79da1bcd77b94c351c
SERVER_TRAFFIC_SECRET_0 997a334266ba1a91eb07ae4e69ea72a7842f77e672a8d7ea33bac783314744 d22de98e1f2a1de1f27a7c0df1fd4ed16c8839997eb083c9e25e02a4b718f1b8
EXPORTER_SECRET 997a334266ba1a91eb07ae4e69ea72a7842f77e672a8d7ea33bac783314744 99c3a8ea9a108a92091652f632baf6d67nb8c538e73ec1ae287cc1f75ac8861d
CLIENT_TRAFFIC_SECRET_0 51f19abff42893a59c0fe7ebcac84b3da2589594a4687d386796c6446356799 1a6cfb3ac436ba73ac92ad589a9087f384827379f833f233479b38046d8d5
SERVER_TRAFFIC_SECRET_0 51f19abff42893a59c0fe7ebcac84b3da2589594a4687d386796c6446356799 171mad411f22312b68380936c0d8f473dd841b86c47ad518498347593f03c
EXPORTER_SECRET 51f19abff42893a59c0fe7ebcac84b3da2589594a4687d386796c6446356799 92c2b8e7a4458a413762f11a3237fae349291179dacc288a98d728c378f88a
CLIENT_RANDOM c3b69cc79b55934f99f1e28f9a9b1a982a50f786426e08ad2712f0cfd15 b6f81aa188d7mb8c63591ee5a748a37d8ee7b74d59978237d435aa7fddfb72766441a75649e
4b8bd84f93a2a543
```

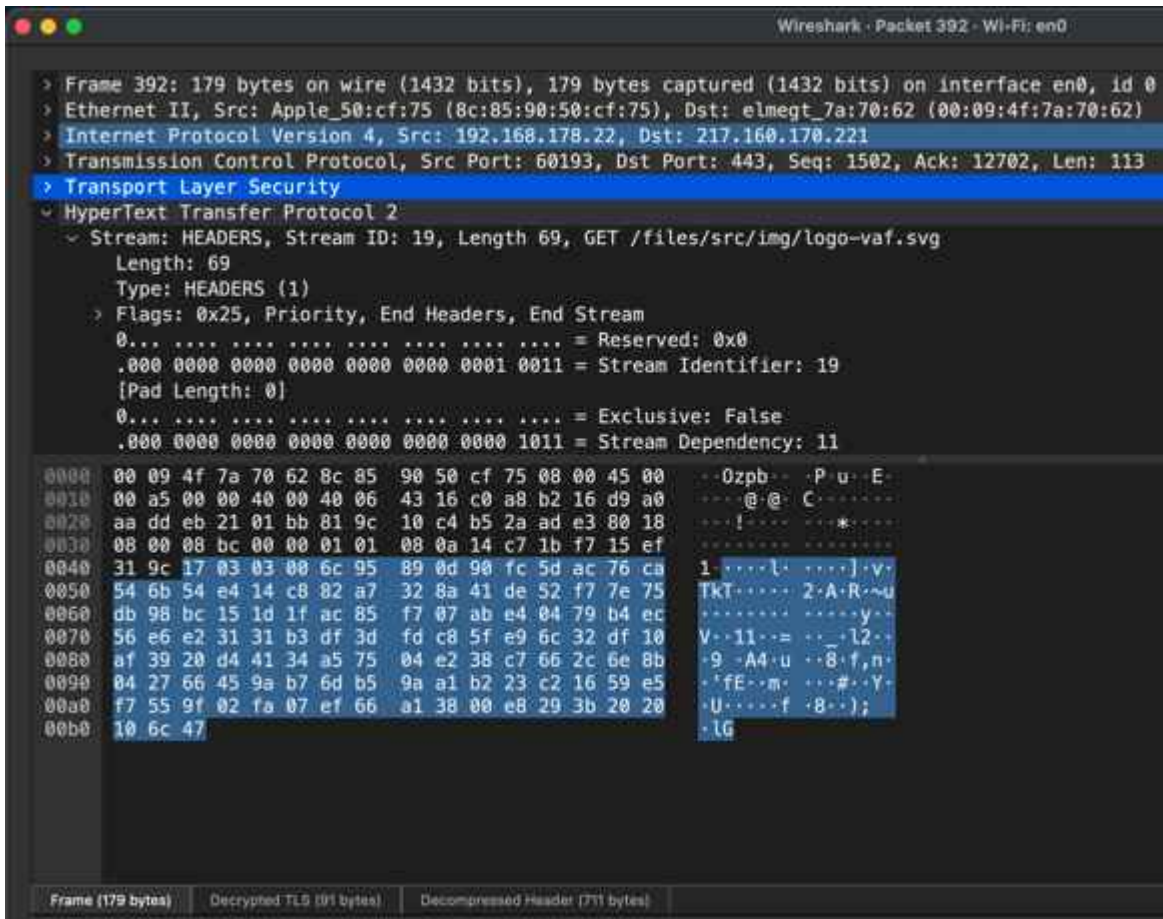
Nach dem Laden der ersten TLS-verschlüsselten Daten zeigt sich am `SSLKEYLOG`, ob die Schlüsseldatei korrekt gespeichert wurde (Abb. 2).

Damit Wireshark die Datei mit den passenden Schlüsseln zum Entschlüsseln heranzieht, ist sie als „(Pre)-Master-Secret log filename“ unter „Preferences/Protocols/TLS“ zu referenzieren (siehe Abbildung 3).



Referenzierung der PMK-Datei in den TLS-Protokolleinstellungen in Wireshark – in diesem Fall /Users/pfister/Desktop/sslkeyfile (Abb. 3).

Sobald der TLS Dissector in Wireshark den Traffic entschlüsselt hat, wird der HTTP2-GET-Request im Klartext lesbar (siehe Abbildung 4). Dass eine Entschlüsselung stattgefunden hat, zeigen die Angabe „HyperText Transport Protocol 2“ unterhalb der Zeile „Transport Layer Security“ und der Hinweis „Decrypted TLS“ im unteren Bereich.



Entschlüsselter HTTP2-GET-Request (Abb. 4).

Wer die Kommandozeile bevorzugt, kann mit tshark arbeiten – es folgt ein Beispiel einer Aufzeichnung und Entschlüsselung per tshark. Zunächst wird wieder die Umgebungsvariable angelegt, gefolgt vom Öffnen des Browsers Mozilla Firefox. Anschließend startet tshark für 60 Sekunden (-a duration:60) ohne direkte Ausgabe (-Q) und schreibt die aufgezeichneten Daten in eine PCAPNG-Datei (-w /Users/pfister/Desktop/tls_decrypt.pcapng). In der letzten Zeile liest tshark die PCAPNG-Datei (-r) mit dem Argument für die Referenz zur Keylog-Datei ein (-o tls.keylog_file:\$SSLKEYLOGFILE) und filtert die Ausgabe über einen Displayfilter auf HTTP (-Y http):

```
export SSLKEYLOGFILE="/Users/<username>/Desktop/sslkeyfile"
open /Applications/Firefox.app
tshark -Q -a duration:60 -w
/Users/pfister/Desktop/tls_decrypt.pcapng &
tshark -r /Users/pfister/Desktop/tls_decrypt.pcapng -o
tls.keylog_file:$SSLKEYLOGFILE -Y http
```


Fazit

Mit der Session-Key-Methode lassen sich selbst aktuelle Protokolle wie TLS 1.3 entschlüsseln. Dafür bedarf es jedoch einer Applikation, die den Sessionschlüssel in eine Logdatei schreibt. Falls dies nicht der Fall ist und Server und Client keine (EC)DHE Cipher Suite nutzen, kann der Analyst als Fallback die RSA-Methode anwenden. Grundsätzlich kann die Möglichkeit einer Entschlüsselung ein Troubleshooting jedenfalls immens erleichtern. Wireshark bietet dafür einen recht einfach zu nutzenden Ansatz. (un@ix.de)

1. Quellen
2. [Weiterführende Informationen finden sich unter ix.de/ztmc.](https://www.ix.de/ztmc)

Spuren kompromittierter E-Mail-Konten analysieren



Spuren kompromittierter E-Mail-Konten analysieren

Nachdem der erste Teil des Playbooks zu gekaperten E-Mail-Accounts zeigte, wie man die Logs mithilfe von Microsofts zentraler Logfunktion einsammelt, erklärt der zweite Teil, wie man sie auf verdächtige Vorgänge auswertet und welche Rückschlüsse sich daraus ziehen lassen.

Nachdem der erste Teil des Playbooks zu gekaperten E-Mail-Accounts zeigte, wie man die Logs mithilfe von Microsofts zentraler Logfunktion einsammelt, erklärt der zweite Teil, wie man sie auf verdächtige Vorgänge auswertet und welche Rückschlüsse sich daraus ziehen lassen.

- Beim ersten Anzeichen verdächtigter Aktivität rund um E-Mail-Accounts sollte man IT-forensische Untersuchungen anstoßen, um zu verstehen, was genau passiert ist. Ausgangspunkt der Analyse sind die gesammelten Logdaten und Artefakte.
- Aussagekräftig im Hinblick auf Eindringlinge ins Firmennetz sind unter anderem fehlgeschlagene Anmeldevorgänge, eingerichtete Mailweiterleitungen oder

neu vergebene Berechtigungen. Solche Hinweise sollten sorgfältig untersucht werden.

- Die Ursachenforschung und eine Nachbereitung sind das A und O nach der Bewältigung von Sicherheitsvorfällen. Daraus abgeleitete technische Maßnahmen sowie die Sensibilisierung von Mitarbeitenden sollen künftige Angriffe zumindest erschweren.

Die umfassendste Datenquelle zur Analyse von Unregelmäßigkeiten oder Verdachtsmomenten für einen Sicherheitsvorfall bietet Microsofts zentrale Logfunktion Unified Audit Log (UAL). Hier werden Benutzer- und Administratoraktivitäten auch unabhängig vom Einsatz zusätzlicher Produkte wie Microsoft Sentinel oder Microsoft Defender for Identity aufgezeichnet (wie die Logdaten im Detail gesichert werden, beschreibt [1]). Die nachfolgenden Schritte zeigen, wie man bei der Analyse vorgeht und die Logdaten sinnvoll durchsuchen kann.

Schritt 4: Untersuchen der Anmeldeaktivitäten

Jedes Mal, wenn sich ein Benutzer bei seinem Konto anmeldet, wird ein Ereignis im UAL erstellt. Dieses Ereignis enthält wichtige Informationen, etwa die Quell-IP-Adresse, die sich unter anderem für eine geografische Suche verwenden lässt. Die Ergebnisse lassen sich mit den erwarteten geografischen Standorten eines Unternehmens und seiner Nutzer vergleichen. Wenn zum Beispiel ein Unternehmen in Deutschland ansässig ist und keine Niederlassung in Asien hat oder das VPN des Unternehmens nicht zu einer IP-Adresse in Asien auflöst, würde man keine Ereignisse aus Asien erwarten. Daher wären Anmeldungen aus Asien in diesem Fall verdächtig.

Natürlich kann es auch sein, dass ein Mitarbeiter sich im Urlaub in Asien befindet und sein Firmenhandy dabei hat, dennoch erfordern diese Ausreißer Aufmerksamkeit. Verdächtige

Anmeldungen kann man durch die Suche nach bestimmten Schlüsselwörtern im UAL entdecken. Neben der IP-Adresse liefern auch die Uhrzeit sowie Informationen zum verwendeten Gerät (UserAgent: Betriebssystem, Browser et cetera) gute Anhaltspunkte. Ob das verwendete Gerät dem Unternehmen bekannt ist und von der IT verwaltet wird oder nicht, lässt sich ebenfalls den Ereignissen entnehmen. Für die Suche nach verdächtigen Anmeldeereignissen kann man folgende Schlüsselwörter verwenden:

Schlüsselwort	Bedeutung des Logeintrags
MailboxLogin	Hinweis auf einen erfolgreichen Log-in-Vorgang
UserLoggedIn	Hinweis auf einen erfolgreichen Log-in-Vorgang
UserLoginFailed	Hinweis auf einen fehlgeschlagenen Log-in-Vorgang
IdsLocked	Hinweis auf einen Brute-Force-Angriff. Der Account wurde gesperrt, da zur viele fehlgeschlagene Anmeldeversuche unternommen wurden.
UserKey="Not Available"	Hinweis auf einen Brute-Force-Angriff. Die Anmeldung ist fehlgeschlagen, da der Benutzeraccount nicht existiert.

Neben Ereignissen rund um das Log-in können auch Fehlermeldungen zur Multi-Faktor-Authentisierung (MFA) Indikatoren für mögliche schädliche Aktivitäten sein. Ein Angreifer könnte das Passwort eines Anwenders ausgespäht haben, um dann an der MFA-Abfrage zu scheitern. UAL-Einträge mit den folgenden Schlüsselwörtern sollten näher untersucht werden:

Schlüsselwort	Bedeutung des Logeintrags
UserStrongAuthClientAuthNRequired	Der Benutzer wird zur Bestätigung einer MFA-Abfrage aufgefordert.
UserStrongAuthClientAuthNRequiredInterrupt	fehlgeschlagene MFA-Abfrage

Schritt 5: Untersuchen von Weiterleitungsregeln

Nachdem ein Angreifer einen Benutzeraccount kompromittiert hat, erstellt er häufig Weiterleitungsregeln, um eingehende E-Mails an ein externes Postfach zu schicken. Auf diese Weise kann er die Aktivitäten eines Opfers kontinuierlich überwachen, ohne sich aktiv in das Konto einzuloggen. Selbst wenn das Passwort eines kompromittierten Kontos zurückgesetzt wird, kann der Angreifer weiterhin E-Mails mitlesen.

Ebenfalls beliebt ist der Einsatz von Weiterleitungsregeln zum automatisierten Löschen von E-Mails, um Spuren, die auf Unregelmäßigkeiten hinweisen, zu verwischen. Auch können Weiterleitungsregeln dazu dienen, Spuren vor dem Anwender zu verstecken, indem E-Mails automatisch als gelesen markiert und in einen anderen Ordner (zum Beispiel in den Junk- oder den RSS-Ordner) verschoben werden.

Einem Angreifer bieten sich in einer Microsoft-365-Umgebung gleich mehrere Möglichkeiten, E-Mails an ein externes Postfach umzuleiten. Er kann zunächst einmal Inbox-Regeln anlegen, um E-Mails auszuleiten. Verfügt das Konto zudem über administrative Berechtigungen, ist auch eine Ausleitung über die globalen Postfacheinstellungen oder Exchange-Transportregeln möglich.

Aktive Inbox-Regeln lassen sich mit der Exchange-Management-

Shell auffinden, falls sie nicht bereits mittels des im ersten Artikel vorgestellten Tools Hawk extrahiert wurden:

```
Get-InboxRule -Mailbox | ? {$_forwardto -or  
$_forwardasattachmentto -or $_redirectto}
```

Auch aktive Mailbox-Weiterleitungen kann die Exchange-Management-Shell anzeigen:

```
Get-Mailbox <identity> | Format-List  
ForwardingSMTPAddress,DeliverToMailboxandForward
```

Der Powershell-Befehl Get-TransportRule liefert eine Übersicht über alle bestehenden Weiterleitungsregeln.

Des Weiteren kann man im UAL potenzielle Angreiferaktivitäten im Zusammenhang mit Weiterleitungsregeln analysieren. Hier lassen sich auch Regeln nachvollziehen, die der Angreifer schon wieder gelöscht hat. Folgende Schlüsselwörter führen zu den relevanten Logeinträgen:

Schlüsselwort	Bedeutung des Logeintrags
New-InboxRule	Anlegen einer neuen Weiterleitungsregel (Inbox-Ebene)
New-TransportRule	Anlegen einer neuen Transportregel (Mail Flow Rule)
Set-Mailbox	Änderungen an den Einstellungen einer Mailbox; kann zum Einrichten einer Weiterleitung auf Mailbox-Ebene verwendet werden
Set-InboxRule	Änderung an einer bestehenden Weiterleitungsregel (Inbox-Ebene)
Set-TransportRule	Änderung an einer bestehenden Transportregel (Mail Flow Rule)
DeliverToMailboxAndForward	Hinweis darauf, dass eine E-Mail an eine andere Mailbox weitergeleitet wurde

Schlüsselwort	Bedeutung des Logeintrags
ForwardingSMTPAddress	Hinweis auf eine erfolgte Weiterleitung einer E-Mail
ForwardingAddress	Hinweis auf eine erfolgte Weiterleitung einer E-Mail
SentTo	Hinweis darauf, wohin eine E-Mail gesendet beziehungsweise weitergeleitet wurde
BlindCopyTo	Hinweis darauf, wohin eine E-Mail gesendet beziehungsweise weitergeleitet wurde
ForwardTo	Hinweis darauf, wohin eine E-Mail gesendet beziehungsweise weitergeleitet wurde

Schritt 6: Persistent Access – Hintertüren entdecken

Im nächsten Schritt gilt es zu prüfen, ob der Angreifer Hintertüren eingerichtet hat. Das würde ihm auch im Fall einer Entdeckung noch Zugriff auf die erbeuteten Konten gewähren. Hier gibt es im Wesentlichen drei beliebte Techniken: App-Kennwörter, das Einrichten schädlicher OAuth-Applikationen und die Manipulation von Berechtigungen.

App-Kennwörter dienen eigentlich der Absicherung von Netzwerkprotokollen, die Microsofts „Modern Authentication“ nicht unterstützen. Um die Sicherheit eines Kontos nicht durch die Verwendung des Kennwortes über ein Protokoll, das nicht dem aktuellen Sicherheitsstand entspricht, zu gefährden, bietet Microsoft die Möglichkeit, ein spezifisches Kennwort einzurichten. Es gilt nur für dieses Protokoll.

Wird es kompromittiert, erhält der Angreifer nur Zugriff zu einem einzelnen Protokoll, zum Beispiel IMAP oder POP, nicht aber zum gesamten Nutzerkonto. Doch Angreifer können diese

Funktion auch missbrauchen, damit sie über ein selbst eingerichtetes App-Kennwort auch nach Änderung des Kennworts im Azure AD noch Zugriff auf die Mails eines Nutzers haben und gegebenenfalls auch weiterhin illegitime Mails verschicken können.

Zur Prüfung auf App-Passwörter sollten Administratoren zum einen im Azure AD die für den jeweiligen Benutzeraccount hinterlegten Authentifizierungsmethoden sichten und zum anderen im Kontext des Kontos selbst die Liste der App-Kennwörter abrufen (siehe ix.de/z2y8).

Anwendungen als Hintertür missbrauchen

Auch Enterprise-Applikationen, die sich mittels OAuth authentifizieren, können als Hintertür zu einem kompromittierten Konto genutzt werden. Berechtigt der Angreifer eine von ihm kontrollierte Enterprise-Applikation zum Zugriff auf das übernommene Konto, erlaubt er damit der Applikation, Aktionen im Kontext des Benutzers durchzuführen.

So ist über diese Applikation auch nach Änderung des Kennworts ein Zugriff mit den gewährten Berechtigungen möglich. Um zu prüfen, ob im Rahmen eines Angriffs Enterprise-Applikationen Berechtigungen erhielten – Microsoft spricht in diesem Zusammenhang von „Illicit Consent Attacks“ –, gibt es mehrere Möglichkeiten.

Administratoren können die Berechtigungen über das Azure-Active-Directory-Portal über den Menüpunkt „Nutzer“ und Auswahl des betroffenen Nutzerkontos prüfen. Eine globale Liste zeigt im Azure AD der Unterpunkt Enterprise-Applikationen. Wer lieber mit PowerShell arbeitet, kann das Skript AzureADPSPermissions.ps1 (siehe ix.de/z2y8) verwenden, um sämtliche OAuth-Berechtigungen eines Tenant in eine CSV-Datei zu exportieren und anschließend zu überprüfen.

Das Hinzufügen von Enterprise-Applikationen beziehungsweise

das Erteilen von Berechtigungen für sie im Analysezeitraum wird im UAL erfasst. Das Werkzeug Hawk extrahiert die Artefakte automatisch (Azure_Application_Audit.csv und Consent_Grant.csv).

Eine Variante zum Phishing mittels OAuth-Applikationen ist das sogenannte Device-Code-Phishing, mit dem sich Office-365-Konten übernehmen lassen. Details zu dieser Angriffstechnik sowie Hinweise zur Detektion und Aufklärung finden sich in einem Artikel des Sicherheitsforschers Nestori Syynimaa (siehe ix.de/z2y8).

Schlüsselwort	Bedeutung des Logeintrags
Add OAuth2PermissionGrant	Einer Enterprise-Applikation wurden Berechtigungen erteilt.
Consent to application	Einer Enterprise-Applikation wurden Berechtigungen durch einen Admin erteilt.
Add app role Assignment grant to use	Ein Benutzer wurde einer Applikation hinzugefügt.

Hat ein Angreifer mehrere Konten eines Unternehmens kompromittiert, kann er sie dazu missbrauchen, Hintertüren einzurichten, indem er den anderen kompromittierten Konten Zugriff auf eine Mailbox gibt. Solange die Verteidiger nicht sämtliche betroffenen Konten identifizieren, behält der Angreifer weiter Zugriff.

Ereignisse im Zusammenhang mit Berechtigungsänderungen lassen sich durch die Suche nach den folgenden Schlüsselwörtern im UAL ausfindig machen:

Schlüsselwort	Bedeutung des Logeintrags
Add-MailboxPermission	Neue Berechtigungen auf ein Postfach wurden vergeben.

Schlüsselwort	Bedeutung des Logeintrags
Add-MailboxFolderPermission	Neue Berechtigungen auf einen Ordner in einem Postfach wurden vergeben.
Add-RecipientPermission	Hinweis darauf, dass einem Benutzer die „Senden als“-Berechtigung zugewiesen wurde.
Set-MailboxFolderPermission	Bestehende Berechtigungen eines Ordners in einem Postfach wurden geändert.

Hat ein Angreifer sogar ein Konto mit administrativen Berechtigungen gekapert, kann er zudem eigene neue Benutzerkonten anlegen, die dann als Hintertür dienen. Auch das hinterlässt Spuren im UAL.

Schlüsselwort	Bedeutung des Logeintrags
Added user	Ein neuer Benutzer wurde angelegt.

Schritt 7: Datenexfiltration analysieren

Bestätigt es sich, dass jemand Unbefugtes Zugriff auf das Unternehmensnetzwerk hatte, stellt sich in erster Linie die Kernfrage: Worauf hat der Angreifer zugegriffen? Dem zugrunde liegt oft die (späte) Erkenntnis über Art und Umfang der Informationen, die mit einem Benutzerkonto prinzipiell erreichbar wären, verbunden mit dem Wunsch, dieses Worst-Case-Szenario irgendwie einzugrenzen.

Hier zunächst die schlechte Nachricht vorweg: Es ist in der Praxis selten möglich, einen Negativbeweis zu führen, also festzustellen, was die Angreifer nicht mitgenommen haben. Die Aussagekraft der Artefakte ist meist begrenzt, da schlicht nicht alles protokolliert wird. In der Regel muss bei einer gesicherten Kompromittierung eines Kontos unterstellt werden, dass der Angreifer alle erreichbaren Inhalte ausgespäht hat. Das hat erhebliche Konsequenzen beispielsweise für die

datenschutzrechtliche Bewertung eines Vorfalls.

Die gute Nachricht ist, dass auch Microsoft das erkannt hat. Konten, die mit einer E5-Lizenz ausgestattet sind, verfügen über eine „erweiterte Überwachung“. Diese Funktion protokolliert unter anderem Zugriffe auf einzelne E-Mails, was die Chance auf den seltenen Negativbeweis zumindest für die Inhalte des Postfachs deutlich verbessert.

Im UAL finden sich dann Einträge der Art MailItemsAccessed. Diese haben unter anderem ein Attribut MailAccessType, das zwischen Bind und Sync unterscheidet.

Operation	Bedeutung des Logeintrags
MailItemsAccessed	Hinweis auf den erfolgten Zugriff auf Inhalte eines Postfachs

Bind-Einträge werden erzeugt, wenn eine einzelne E-Mail abgerufen wird. Die ID der Nachricht steht dann im Attribut InternetMessageId. Die Protokollierung unterliegt jedoch einer wichtigen Einschränkung: Werden innerhalb von 24 Stunden mehr als 1000 Zugriffe dokumentiert, wird die Protokollierung für Bind-Ereignisse für 24 Stunden ausgesetzt (Throttle).

Zuerst sollte also geprüft werden, ob das UAL Einträge des Typs MailItemsAccessed für die zu untersuchende Mailbox enthält. Anschließend gilt es auszuschließen, dass ein Throttling stattgefunden hat. Dazu schaut man, ob es bei den MailItemsAccessed-Ereignissen welche gibt, die beim Attribut IsThrottled den Wert True vermerkt haben. Im Idealfall gibt es keinen solchen Eintrag.

Welche Sitzung gehört zu wem?

Der nächste Schritt besteht darin, die zum Angreifer gehörenden Sitzungen zu ermitteln. Dafür gleicht man die MailItemsAccessed-Vorgänge im UAL mit den Informationen des Angreifers (verdächtige Log-in-Aktivitäten, IP-Adressen, Zeitstempel, Art des Zugriffs) und den Informationen über den

legitimen Anwender ab. Die Einträge haben mitunter mehrere Session-IDs und IP-Adressen für ein Benutzerkonto. Anhand der in den vorangegangenen Schritten ermittelten Kompromittierungsindikatoren lässt sich feststellen, welche Sitzungen wahrscheinlich legitim oder gültig sind. Einige Sitzungen haben möglicherweise keine Session-ID, weil für die Anmeldung eine alte (Legacy-)Authentifizierung verwendet wurde. Die verdächtigen MailItemsAccessed-Einträge werden dann weiter analysiert.

Sync-Einträge entstehen immer dann, wenn ein E-Mail-Client, beispielsweise Outlook, ein Postfach synchronisiert und dabei Inhalte auf einen lokalen Computer herunterlädt. Hierbei entsteht kein Logeintrag pro Element, sondern pro Ordner des Postfachs. Finden sich im UAL MailItemsAccessed-Einträge mit dem MailAccessType Sync, die dem Angreifer zugeordnet werden, so muss man davon ausgehen, dass alle E-Mails im synchronisierten Ordner kompromittiert wurden.

Zuletzt bleiben die Bind-Vorgänge, die dem Angreifer zugeordnet werden. Diese enthalten eine InternetMessageID. Um damit auf die eigentlichen Nachrichten schließen zu können, ist es notwendig, das Message Trace Log mit den IDs abzugleichen. Leider reicht das Message Trace Log nicht so weit zurück wie die Einträge im UAL, sondern lediglich zehn Tage. Auch lässt sich die InternetMessageID nicht als Suchparameter im Rahmen einer Suche nach Beweismitteln (E-Discovery) verwenden.

Können E-Mails nicht mehr über das Message Trace Log zugeordnet werden, bleibt lediglich der Weg, das Postfach selbst zu exportieren und die E-Mails zu durchsuchen. Die ID ist in den Eigenschaften der E-Mails gespeichert. Der Export des Postfachs lässt sich außerdem über die E-Discovery-Funktion realisieren, die auch bereits gelöschte Elemente berücksichtigt (sofern entsprechende Aufbewahrungsrichtlinien konfiguriert sind und die Elemente noch vorgehalten werden).

Rekonstruieren, was geklaut wurde

Wie beschrieben können E-Mails auch über Weiterleitungsregeln abgegriffen werden. Findet man bei einer Untersuchung solche Regeln, kann sowohl das UAL (siehe Schritt 5) wie auch die Logik der Regeln selbst Aufschluss über die betroffenen Inhalte geben. Neben dem Abgleich der Einträge im UAL mit dem Message Trace Log sollte die Mailbox nach den Parametern der Regel(n) durchsucht werden.

Sofern ein Angreifer Zugang zu einem Konto mit administrativen Berechtigungen und der E-Discovery-Suche hatte, kann er auch auf diesem Weg Inhalte gesucht und exportiert haben. Hinweise darauf lassen sich wieder im UAL finden.

Analog zu den E-Mails sind alle weiteren Inhalte zu berücksichtigen, die mit dem kompromittierten Konto für den Angreifer erreichbar waren. Das beinhaltet sowohl in OneDrive geteilte Dateien wie Teams-Nachrichten und SharePoint-Seiten als auch sämtliche nachgelagerten Applikationen, die Azure AD zur Authentifizierung verwenden. Die Analyse ist allerdings oft sehr individuell und würde den Rahmen dieses Artikels sprengen.

Schritt 8: Remediation

Nachdem die Aktivitäten eines Angreifers nachvollzogen wurden, gilt es, alles rückgängig zu machen, also alle gefundenen Weiterleitungsregeln, Enterprise-Applikationen, App-Kennwörter et cetera zu entfernen und die Kennwörter der betroffenen Konten, falls noch nicht geschehen, zurückzusetzen. Auch sollten alle Analysen und eingeleiteten Maßnahmen dokumentiert und mit den zugehörigen Logdateien aufbewahrt werden.

Zeigte die Untersuchung einen unberechtigten Zugriff auf Postfächer, handelt es sich um einen meldepflichtigen Vorfall gemäß der DSGVO. Dementsprechend ist eine Erklärung an den zuständigen Landesdatenschutzbeauftragten verpflichtend. Dabei

gilt es, die gesetzlichen Fristen zu beachten. Binnen 72 Stunden ab dem Zeitpunkt der Kenntnisnahme muss die Meldung erfolgen. Zu diesem Zeitpunkt ist gegebenenfalls noch nicht das gesamte Ausmaß des Vorfalls bekannt. In diesem Fall sollte die Meldung einfach alle bisher gesicherten Informationen enthalten. Die Meldung sollte durch den benannten Datenschutzbeauftragten des betroffenen Unternehmens erfolgen.

Neben den Datenschutzbehörden müssen gegebenenfalls auch die betroffenen Personen informiert werden. Dies ist dann der Fall, wenn besonders heikle personenbezogene Daten gemäß Art 9 DSGVO – also beispielsweise religiöse oder weltanschauliche Überzeugungen oder Gesundheitsdaten – betroffen sind. In diesem Fall sind die betroffenen Personen direkt zu benachrichtigen. Die Prüfung einer solchen Meldepflicht obliegt dem Datenschutzbeauftragten. Gegebenenfalls sollte bei Verdacht auf einen solchen Fall juristischer Beistand hinzugezogen werden.

Schritt 9: Root Cause Analysis – woran liegt's?

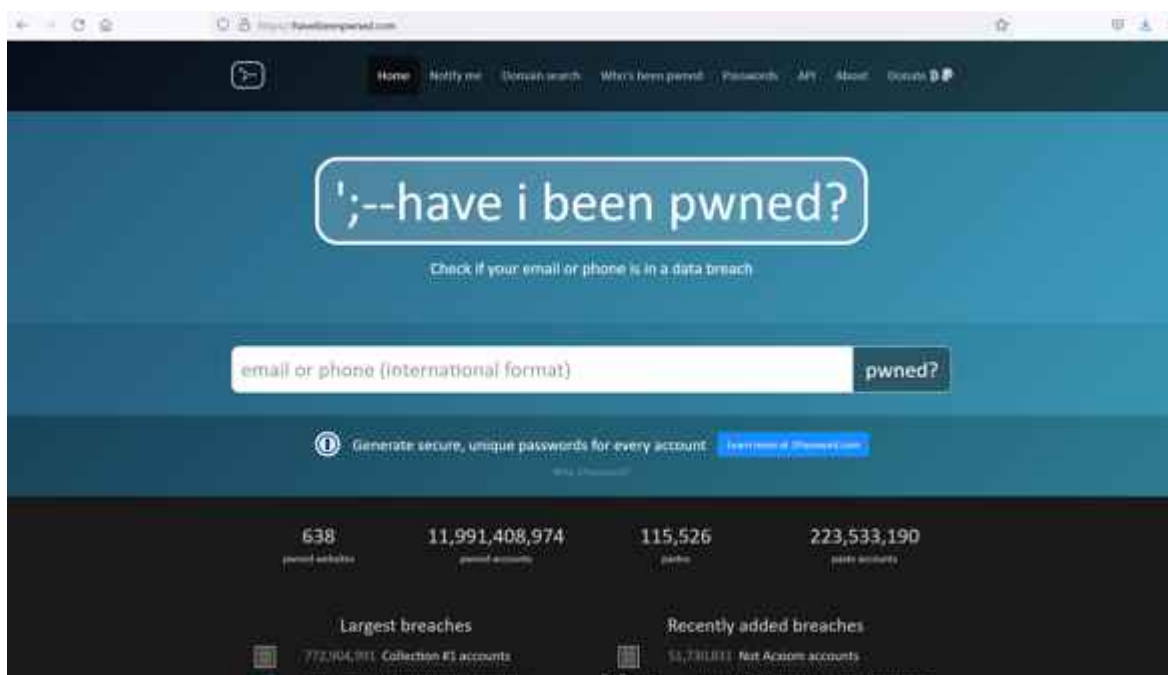
Nachdem aufgeklärt ist, wie ein Angreifer vorgegangen ist und was er genau getan hat, bleibt noch die Frage, wie das passieren konnte. Wie hat er initial Zugang erhalten?

Auch hier ist leider keine pauschale Anleitung möglich, doch die häufigsten Ursachen sind folgende:

- Password Spraying / Brute Force / einfach zu erratende Passwörter: Allen drei Szenarien ist gemeinsam, dass sie in der Regel mit mehrfachem Ausprobieren einhergehen. In den Logs äußert sich dies durch multiple fehlgeschlagene Log-in-Versuche bei einem oder mehreren Konten, ausgehend von derselben IP-Adresse und/oder ähnlichen Parametern wie User-Agent, Protokoll und Zeitpunkt.
- (Spear-)Phishing: Bei einem Phishingangriff erhält das

Opfer eine E-Mail, die einen Link oder einen Anhang enthält, über den die Zugangsdaten abgegriffen werden (funktioniert teilweise auch bei MFA) oder eine Enterprise App via OAuth-Berechtigungsanfrage untergeschoben wird. In dem Fall sind keine gehäuften fehlgeschlagenen Log-in-Versuche zu beobachten. Stattdessen gilt es, die Phishingmail im Postfach oder den aufgerufenen Link ausfindig zu machen.

- Password Re-use / Leaked Credentials: Oft verwenden Anwender ein Passwort für mehrere Dienste und Konten oder recyceln ein privates Passwort für Firmenzwecke. In dem Fall kann es sein, dass das Kennwort bei einem der anderen Dienste ausgespäht wurde und dann für die Anmeldung am Microsoft-365-Account ausprobiert wird. Auch hier ist nicht unbedingt eine gehäufte Anzahl an Fehlversuchen zu beobachten, sofern nicht zusätzlich MFA aktiviert ist. Um der Ursache in dem Fall näherzukommen, empfiehlt es sich, mit dem Benutzer ein offenes Gespräch zu führen oder die Unternehmens-E-Mail des Anwenders bei seriösen Diensten wie haveibeenpwned.com einzugeben (siehe Abbildung).



Ob ein Passwort geleakt wurde, kann man beispielsweise bei Diensten wie „Have I Been Pwned“ herausfinden. Dieser Dienst

des australischen Sicherheitsforschers Troy Hunt hat einen guten Ruf, da er nicht das Passwort selbst, sondern nur den Benutzernamen abfragt.

Nach der erfolgreichen Bewältigung des potenziellen oder realen Sicherheitsvorfalls sollte immer auch geprüft werden, welche Lektionen man daraus lernen kann und welche Maßnahmen zu ergreifen sind, damit ähnliche Vorfälle in Zukunft seltener oder gar nicht mehr vorkommen. Dabei soll es explizit keine Schuldzuweisungen geben, das Stichwort lautet hier vielmehr „Blameless Post Mortem“.

Awareness-Maßnahmen und Schulungen können gängige Betrugsmuster vermitteln und damit die Anfälligkeit der Mitarbeitenden für solche Angriffe verringern. Klar definierte Prozesse zur Veranlassung von Zahlungen helfen außerdem, bestimmte Arten von finanziellem Betrug zu erschweren. Häufig werden aber im Rahmen der Vorfallsbehandlung vor allem technische Gegebenheiten identifiziert, die die Kompromittierung erleichtert oder die Untersuchung des Vorfalls erschwert haben. So ist es hilfreich, die SPF-, DKIM- oder DMARC-Konfiguration (Sender Policy Framework; DomainKeys Identified Mail; Domain-based Message Authentication, Reporting and Conformance) nachzurüsten, falls sie im Vorfeld des Vorfalls noch nicht aktiv war, die Protokollierung lässt sich verbessern, wenn Logs für die Aufklärung des Angriffs fehlten, oder das Installieren von OAuth-Anwendungen kann für Nutzer des Tenants eingeschränkt werden, falls Angreifer solche Anwendungen als Hintertür installiert haben.

Microsoft gibt im Rahmen einer Referenzarchitektur zahlreiche Hinweise für das Absichern von Microsoft-365- und Azure-AD-Umgebungen (siehe ix.de/z2y8), die im Nachgang eines Vorfalls (re-)evaluiert werden und bei Bedarf in das Sicherheitskonzept des Unternehmens integriert werden können. Dedizierte Dienste wie Microsoft Defender for Office, Microsoft Defender for Identity oder Microsoft Defender for Cloud Apps können gegen Angriffe schützen oder bei ihrer Entdeckung und Aufbereitung helfen. Allerdings sind sie häufig nur in den teureren

Lizenzen der Microsoft-Produkte enthalten oder müssen sogar separat lizenziert werden. (ur@ix.de)

1. Quellen
2. [Jens Lüttgens, Dominik Oepen; E-Mail-Betrug in MS-365-Umgebungen; iX 12/2022, S. 102](#)
3. [Vertiefende Microsoft-Artikel, das erwähnte PowerShell-Skript sowie die Microsoft-Referenzarchitektur sind über \[ix.de/z2y8\]\(https://ix.de/z2y8\) zu finden.](#)



Introducing a new phishing technique for compromising Office 365 accounts

The ongoing global phishing campaigns againsts Microsoft 365 have used various phishing techniques.

Currently attackers are utilising forged login sites and OAuth app consents. In this blog, I'll introduce a new phishing technique based on Azure AD device code authentication flow. I'll also provide...

IT-Recht 2023: Viele neue EU-Regeln



IT-Recht 2023: Viele neue EU-Regeln

Im kommenden Jahr muss sich die IT-Welt mit etlichen neuen rechtlichen Regulierungen auseinandersetzen, viele davon auf EU-Ebene. Unter anderem sollen die Internetgiganten stärker an die Leine genommen werden. Aber auch kleine Unternehmen müssen handeln.

Im kommenden Jahr muss sich die IT-Welt mit etlichen neuen rechtlichen Regulierungen auseinandersetzen, viele davon auf EU-Ebene. Unter anderem sollen die Internetgiganten stärker an die Leine genommen werden. Aber auch kleine Unternehmen müssen handeln.

Es gibt einen Grund, warum auf EU-Ebene derzeit viele Gesetzgebungsvorhaben im IT-Bereich forciert werden: die im

Frühjahr 2024 anstehende Europawahl. Insbesondere die EU-Kommission möchte bis dahin möglichst alle ihre in der Agenda „Priorities 2019 – 2024 – A Europe fit for the digital age“ gesetzten Ziele erreichen. Die Amtszeit der derzeitigen Kommission endet mit der Legislaturperiode des Europäischen Parlaments. Anschließend wird eine neue EU-Kommission gebildet, die sich dann eine neue IT-Rechts-Agenda geben dürfte.

2023 werden zunächst zahlreiche EU-Gesetze in Kraft treten, die bereits im Jahr 2022 beschlossen wurden. Hierzu zählt der **Digital Markets Act (DMA)**, der am 1. November 2022 in Kraft getreten und ab dem 2. Mai 2023 wirksam ist. Er sieht vor, dass es auf Plattformen der Gatekeeper im Internet fair zugeht, wie es auf einer Webseite der EU-Kommission heißt. Anhand objektiver Kriterien wird festgestellt, ob es sich bei einer Onlineplattform um einen solchen Gatekeeper handelt. Relevant sind dabei insbesondere die wirtschaftliche Position und die Nutzerzahlen.

Der DMA sieht vor, dass Gatekeeper künftig diskriminierungsfrei ihre Plattformen für den Absatz von Waren und Dienstleistungen durch Dritte zur Verfügung stellen müssen. Dies gilt auch für die dabei von Nutzern auf der Plattform hinterlassenen Daten. Eigene Waren und Dienstleistungen darf der Gatekeeper dabei nicht bevorzugen, auch darf er Nutzer nicht vom Deinstallieren von Apps abhalten. Außerhalb der Plattform darf er Nutzer nicht ohne deren Einwilligung bewerben. Die Bußgelder können bis zu 20 Prozent des weltweiten Jahresumsatzes betragen.

Länderübergreifende Dienste

Beim **Digital Services Act (DSA)** hat sich die EU auf eine längere Frist zwischen dem Inkrafttreten am 16. November 2022 und dem Wirksamwerden am 17. Februar 2024 verständigt. Hintergrund hierfür sind die zahlreichen und teils tiefgreifenden Vorgaben für sehr viele Unternehmen, die

Leistungen rund um das oder im Internet anbieten. Im Wesentlichen geht es bei der Regulierung darum, Verbraucher und ihre Grundrechte besser zu schützen, einen einheitlichen Rechtsrahmen zu schaffen und – vor allem auch für kleinere Serviceanbieter, KMU oder Start-ups – den Zugang zu EU-weiten Märkten zu vereinfachen. Nicht zuletzt liegt ein Schwerpunkt des DSA auf der Minderung systemimmanenter Risiken wie Manipulation oder Desinformation (siehe [ix.de/zqe9](https://www.ix.de/zqe9)).

Neben den üblichen Folgen bei Rechtsverstößen wie wettbewerbsrechtlichen Abmahnungen, einstweiligen Verfügungen und dergleichen sieht der DSA Bußgelder von bis zu sechs Prozent des weltweiten Jahresumsatzes des Anbieters vor. Betroffen vom DSA sind „vermittelnde Online-Dienste“. Hierzu zählen Vermittlungsdienste mit einem eigenen Infrastrukturnetz, etwa Internetanbieter, DNS-Registrierstellen und Hosting-Dienste im Bereich Cloud und Webhosting. Erfasst sind des Weiteren Onlineplattformen wie Onlinemarktplätze, App-Stores oder Social-Media-Plattformen. Der DSA sieht in den Regelungen zum Anwendungsbereich keine Ausnahmen für nicht kommerzielle Anbieter vor. Also dürften Mastodon und gegebenenfalls auch Wikipedia unter den Anwendungsbereich fallen.

Die betroffenen Unternehmen sind gut beraten, das Jahr 2023 zur Vorbereitung zu nutzen. Es gilt, die Compliance mit dem DSA zu schaffen, die AGB anzupassen und womöglich auch die angebotenen Leistungen selbst [1].

Der DSA wird in Fachkreisen auch als „Biest“ bezeichnet, denn die Vorgaben sind sehr weitreichend. Neben Tech-Giganten dürften beispielsweise auch einzelne geschäftliche WLAN-Betreiber betroffen sein. Mit Abmahnungen bei DSA-Verstößen ist ab Februar 2024 zu rechnen. Diese Abmahnwelle könnte deutlich größere Ausmaße annehmen als die derzeitige bei der Verwendung dynamischer Google-Fonts.

Kryptoregulierung verspätet sich

Eigentlich sollte die Verordnung **Markets in Crypto-Assets (MiCA)** bereits 2022 verabschiedet werden und in Kraft treten. Überraschend vertagte das EU-Parlament die Beschlussfassung jedoch auf 2023. Inhaltlich bestand weitgehend Einigkeit zwischen EU-Rat, -Kommission und -Parlament. MiCA regelt die „digitale Darstellung eines Wertes oder eines Rechts, das elektronisch transferiert und gespeichert werden kann“, wenn dafür „die Distributed-Ledger-Technologie oder eine vergleichbare Technologie verwendet“ wird. Non-Fungible Tokens (NFT) sind nach derzeitigem Stand als Ergebnis längerer Diskussionen auf Gesetzgebungsebene nicht von der Verordnung betroffen. Die Verordnung ist Teil des EU-Pakets zur Digitalisierung des Finanzwesens.

Die MiCA-Verordnung soll EU-weit Krypto-Assets regulieren. Sie nimmt Emittenten und Dienstleister in den Fokus. Neben dem Anlegerschutz durch Transparenz- und Offenlegungspflichten stehen unter anderem die Verhinderung von Marktmissbrauch und Geldwäsche im Raum. Für zahlreiche Dienstleistungen wird zukünftig die Erlaubnis der Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) erforderlich sein. Die Anforderungen ähneln denen an Finanzinstitute.

Kryptodienstleister müssen ihren Sitz und mindestens einen Geschäftsleiter in der EU haben. Sie müssen die BaFin über das Unternehmen sowie dessen Gesellschafter und Geschäftsleiter umfassend informieren. Die Geschäftsleiter müssen zudem fachlich geeignet und zuverlässig, die Geschäftsorganisation muss ordnungsgemäß und angemessen sein. Maßnahmen gegen Geldwäsche und die ausreichende Organisation der Compliance sind ebenso vorgeschrieben wie ein professionelles Beschwerdemanagement und die Pflicht, eigene Vermögenswerte von denen der Kunden zu trennen.

Sichere Standards für vernetzte Produkte

Am 15. September 2022 hat die EU-Kommission einen ersten Entwurf für einen **Cyber Resilience Act (CRA)** vorgestellt, der nun durch das Gesetzgebungsverfahren und die Abstimmungen zwischen EU-Kommission, -Rat und -Parlament läuft. Das Gesetz soll gemeinsame Cybersicherheitsstandards für vernetzte Geräte und Dienste („Produkte mit digitalen Anteilen“) festlegen und damit spürbar zur Bekämpfung von Cyberkriminalität beitragen. Mit seiner Verabschiedung ist 2023 zu rechnen, 24 Monate nach Inkrafttreten wird es wirksam. Auf Hersteller solcher Produkte kommt aber bereits nach 12 Monaten eine Berichtspflicht zu, wenn in einem Produkt mit digitalen Elementen eine aktiv ausgenutzte Sicherheitslücke auftritt.

Die geplanten Regelungen reichen von der Pflicht von Herstellern und Dienstleistern, ein angemessenes Niveau an Cybersicherheit einzuhalten, bis hin zum Verkaufsverbot für Produkte mit bekannten Schwachstellen. Produkte sollen nur noch in Verkehr gebracht werden, wenn sie im Sinne von Security by Default konfiguriert sind. Zudem müssen Angriffsflächen und mögliche Auswirkungen von Attacken systemseitig begrenzt sein.

Für kritische Produkte sollen zwei Kategorien eingeführt werden. Die Anforderungen an die Compliance mit den CRA-Vorgaben sollen für Hersteller von Desktop- und Mobilgeräten, virtualisierten Betriebssystemen, Ausstellern digitaler Zertifikate, Allzweck-Mikroprozessoren, Kartenlesegeräten, Robotersensoren, intelligenten Zählern und IoT-Geräten jeglicher Art, Routern und Firewalls für den industriellen Einsatz deutlich höher sein als für andere Produkte mit digitalen Inhalten. Der CRA-Entwurf sieht Bußgelder bis 15 Millionen Euro beziehungsweise 2,5 Prozent des weltweiten Jahresumsatzes vor. In ersten Stellungnahmen warnen Branchenvertreter davor, kleine und mittlere Unternehmen durch allzu hohe und kostspielige Sicherheitsanforderungen vom Markt

auszuschließen.

Auf Finanzunternehmen kommen bereits 2023 im Bereich Cybersicherheit zahlreiche Hausaufgaben zu. Am 10. November 2022 hat das EU-Parlament den **Digital Operational Resilience Act (DORA)** verabschiedet. Ziel ist es, bestehende Standards für die Cybersicherheit zu vereinheitlichen. Das soll die digitale Betriebsstabilität von EU-Finanzunternehmen gewährleisten. Geplant ist ein detailliertes und umfassendes Rahmenwerk. DORA soll nach einer Umsetzungsfrist von zwei Jahren wirksam werden. Die Vorgaben gelten damit zum Jahreswechsel 2024/2025 (zu DORA siehe separaten Artikel ab [Seite 92](#)).

Lange erwartet: die NIS2-Richtlinie

Knapp zwei Jahre nach dem Kommissionsvorschlag hat ebenfalls im November 2022 das EU-Parlament der NIS2-Richtlinie zugestimmt. Die noch ausstehende Zustimmung durch die EU-Staaten gilt in Fachkreisen als Formsache. **NIS2** steht für die überarbeitete zweite Fassung der 2016 verabschiedeten **Directive on Security of Network and Information Systems**. Richtlinien sind anders als Verordnungen oder Acts durch die EU-Mitgliedsstaaten in nationales Recht umzusetzen. Ihr Ziel ist die Gewährleistung eines hohen gemeinsamen Sicherheitsniveaus von Netz- und Informationssystemen in der Union.

Geplant ist, durch NIS2 den Anwendungsbereich der bisherigen NIS1-Richtlinie drastisch auszuweiten. Unternehmen mit mehr als 50 Mitarbeitern und einem Jahresumsatz von mehr als 10 Millionen Euro sollen künftig unter NIS2 fallen, wenn sie in einem kritischen Sektor tätig sind. Auch die Auflistung, was als kritischer Sektor einzustufen ist, soll signifikant erweitert werden. Danach fallen künftig etwa auch Hersteller von Medizingeräten, Labore, Cloud-Provider, Rechenzentren und Content-Delivery-Netzwerke darunter. Zum etwas schwächer regulierten „wichtigen Sektor“ zählen künftig der gesamte

industrielle Sektor, Hersteller von Computern sowie die Branchen Maschinenbau und Mobility.



Die von vielen lange ersehnte NIS2-Richtlinie weitet den Geltungsbereich ihres Vorgängers erheblich aus. Zahlreiche weitere Branchen gelten nun als „kritischer Sektor“.

Betroffene Unternehmen müssen Risikoanalyse- und Sicherheitskonzepte für die Informationssysteme, die Bewältigung von Zwischenfällen, die Offenlegung von Schwachstellen sowie die Gewährleistung der Sicherheit in der Lieferkette schaffen. Die Aufsichtsmaßnahmen und Durchsetzungsanforderungen der nationalen Behörden sollen strenger gefasst werden. Der Bußgeldrahmen soll 10 Millionen Euro oder bis zu zwei Prozent des weltweiten Jahresumsatzes umfassen.

Binnen 18 Monaten nach Inkrafttreten sollen die Mitgliedsstaaten die NIS2-Richtlinie umgesetzt haben. Betroffene Unternehmen müssen sich also auf erheblich verschärfte Vorgaben in puncto Cybersicherheit ab 2024 oder spätestens 2025 einstellen. Angesichts des Mangels an Fachkräften in diesem Bereich und des benötigten Vorlaufs für eine Compliance mit den NIS2-Vorgaben müssen sich die Verantwortlichen in Unternehmen spätestens ab 2023 mit der konkreten Umsetzung beschäftigen. Auf Betreiber kritischer

Infrastrukturen kommt am 1. Mai 2023 auf jeden Fall eine bereits beschlossene Pflicht nach dem BSI-Gesetz zu. Sie sind dann verpflichtet, Systeme zur Angriffserkennung zu verwenden.

Ein weiteres Großprojekt der EU ist der **Artificial Intelligence Act (AI Act)**. Nachdem die EU-Kommission bereits im April 2021 einen ersten Gesetzentwurf vorgelegt hat, fand erst im Oktober 2022 die erste Plenarsitzung des EU-Parlaments dazu statt. Ein Grund für die lange Dauer des Verfahrens dürften die über 3000 Änderungsvorschläge sein, mit denen sich das Parlament bei der Regulierung des Einsatzes von künstlicher Intelligenz befassen muss. Die EU beabsichtigt mit dem AI Act einen einheitlichen Rechtsrahmen für vertrauenswürdige KI-Systeme zu schaffen sowie einheitliche Regeln für die Entwicklung, Vermarktung und Verwendung von KI innerhalb der EU im Einklang mit ihren Werten und den Grundrechten.

Schwieriges Ringen um Kompromisse

In Details ist der AI Act sehr umstritten. Der Anwendungsbereich, aber auch der Einsatz biometrischer Erkennungssysteme und ihr potenzieller Missbrauch stehen neben anderen Aspekten im Mittelpunkt der Diskussion. Ein Kompromissvorschlag sieht vor, Behörden in Drittstaaten vom AI Act auszunehmen, wenn sie künstliche Intelligenz im Rahmen von Vereinbarungen über internationale oder justizielle Zusammenarbeit verwenden und ein Angemessenheitsbeschluss der EU-Kommission nach der DSGVO vorliegt. Ausnahmen wird es sicher für die militärische Nutzung und womöglich auch für Forschung und Entwicklung geben. Der EU-Rat fordert zudem eine Beschränkung des Anwendungsbereichs auf maschinelles Lernen.

Angst vor kollektiver biometrischer Überwachung

Strittig ist, welche Ausnahmen es für das pauschale Verbot von

Echtzeit-Fernererkennungssystemen zur biometrischen Identifizierung von Personen im öffentlichen Raum geben soll. Einige EU-Parlamentarier haben Sorge, dass die Zulassung der Identifizierung von Entführungsoptionen und Kriminellen sowie zur Abwehr von unmittelbar drohenden Terroranschlägen zur Überwachung der Gesellschaft quasi durch die Hintertür führen kann. Vereinzelte Forderungen sehen vor, das Verbot auch auf den privaten Bereich auszudehnen und auch durch Streichung des „Echtzeit-Erfordernisses“ eine nachträgliche Identifizierung zu untersagen.

Der AI Act wird einen risikobasierten Regelungsansatz verfolgen. KI-Systeme sollen in die vier Kategorien minimales, geringes, hohes oder unannehmbares Risiko eingestuft werden. Im unteren Bereich stehen Transparenzanforderungen und sektorale Regulierungen im Raum. Erfasst werden beispielsweise Systeme, die mit Menschen interagieren oder Emotionen anhand biometrischer Daten erkennen, sowie Systeme, die Inhalte erzeugen oder manipulieren. Unter Letzteres würden auch Deepfakes, also realistisch wirkende Medieninhalte fallen, die durch KI-Systeme geändert oder verfälscht wurden.

Für KI-Systeme mit hohem Risiko sind hohe Anforderungen an das Risikomanagement, die Datenqualität und die technische Dokumentation vorgesehen. Eine hochrangige Expertengruppe soll hierfür Mindestanforderungen gemäß definierten Ethik-Leitlinien festlegen. Diskutiert wird darüber hinaus eine Konformitätsbewertung, die vor Einsatz des betreffenden KI-Systems positiv ausfallen muss.

Als unannehmbar riskante KI-Systeme werden die genannten biometrischen Systeme zur Fernidentifizierung, aber auch Social Scoring durch Behörden (wie bereits in China praktiziert) sowie manipulative Systeme mittels Techniken der unterschwelligen Beeinflussung Schutzbedürftiger eingestuft. Für sie ist ein generelles Verbot vorgesehen. Verstöße gegen den AI Act sollen durch beträchtliche Bußgelder geahndet werden. Diskutiert wird über einen Rahmen von bis zu 30

Millionen Euro oder sechs Prozent des weltweiten Jahresumsatzes.

US-EU-Datenschutz, die Dritte!

Was noch? Spannend wird sein, ob die EU-Kommission aller Kritik zum Trotz im Frühjahr 2023 einen sogenannten Angemessenheitsbeschluss gemäß Artikel 45 der Datenschutz-Grundverordnung fassen wird, der dem Datenschutz in den USA „ein angemessenes Schutzniveau“ bescheinigt. Seit der Europäische Gerichtshof in seinem viel beachteten Schrems-II-Urteil den **EU-US Privacy Shield** kassiert hat, ist die Übermittlung personenbezogener Daten aus der EU in die USA deutlich erschwert.

Im Oktober 2022 hatte US-Präsident Biden eine Executive Order unterzeichnet, mit der ein angemessenes Datenschutzniveau aus EU-Sicht geschaffen werden soll. Zahlreiche Datenschützer wie der scheidende Landesdatenschutzbeauftragte Baden-Württembergs Stefan Brink, aber auch der Datenschutzaktivist Max Schrems zweifeln daran, dass die Executive Order ausreicht. Der EuGH dürfte erneut mit der Rechtslage befasst werden. Ein Ende der Gemengelage ist nicht absehbar.

Ungeachtet dessen dürften die von Unternehmen getroffenen Maßnahmen und Verträge auch weiterhin nicht den Bestimmungen der DSGVO entsprechen. Seit Ende 2022 gelten neue Vorgaben für die Standardvertragsklauseln. Sie sind derzeit eine der wenigen Möglichkeiten, den Datentransfer in die USA rechtskonform auszugestalten. Die Datenschutzbehörden dürften 2023 mit einer Durchsetzung der Änderungen beginnen und gegebenenfalls signifikante Bußgelder verhängen.

Um die in den letzten Jahren heftig diskutierte **E-Privacy-Verordnung** ist es zuletzt sehr ruhig geworden. Sie soll die DSGVO ergänzen und weiter gehende Rahmenbedingungen für den Umgang mit personenbezogenen Daten im Bereich der elektronischen Kommunikation schaffen. In erster Linie soll es

Regelungen etwa zu Cookies oder Trackern geben. Diskutiert werden auch Vorgaben für Direktmarketing und Teilnehmerverzeichnisse. Ob die Verordnung nun endlich 2023 das Licht der Welt erblicken wird, ist allerdings mehr als fraglich. Aber selbst wenn, dürfte sie nicht vor 2025 wirksam werden.

Weniger wegwerfen, mehr reparieren

Mitte November 2022 haben sich die EU-Mitgliedsstaaten und die EU-Kommission auf neue **Ecodesign-Vorgaben** geeinigt. Sie sollen 2023 formal verabschiedet und nach einer Umsetzungsfrist von 21 Monaten wirksam werden. Eingeführt werden soll ein **Recht auf Reparatur**. Hersteller von Smartphones, Tablets und Co. müssen danach Reparaturanleitungen und für die Dauer von sieben Jahren bestimmte Ersatzteile wie Displays und Batterien verfügbar halten. Software-Updates müssen fünf Jahre lang bereitgestellt werden. Sie dürfen die Geräteperformance nicht beeinträchtigen. Schließlich sollen die Rechte von Dienstleistern gestärkt werden, die Gerätereparaturen anbieten.

2023 dürfte die Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) eine neue Fassung ihres Rundschreibens **Mindestanforderungen an das Risikomanagement (MaRisk)** veröffentlichen. Es wird die derzeit gültige Fassung dieses Rundschreibens vom August 2021 ersetzen. Aus IT-Sicht interessant sind die Diskussionen rund um IT-Sicherheit und IT-Zugang zu Handelsplattformen aus dem Homeoffice. Infolge der Coronapandemie haben zahlreiche Finanzdienstleister gefordert, den strengen Ansatz aufzuweichen, dass beispielsweise ihr Aktienhandel nur „in Geschäftsräumen“ stattfinden darf. Letztlich haben Änderungen in der MaRisk zahlreiche Auswirkungen auf die im Finanzwesen eingesetzten IT-Systeme. Relevant ist hier auch das 2021 überarbeitete Rundschreiben **Bankaufsichtsrechtliche Anforderungen an die IT**, kurz **BAIT**, das die MaRisk konkretisiert. Womöglich steht auch

dieses 2023 zur Überarbeitung an.

Weitergehen dürfte es 2023 auch mit den Vorbereitungen für einen **European Chips Act**, der die Wettbewerbsfähigkeit und Resilienz der Chipindustrie in der EU signifikant stärken soll. Am 24. September 2023 wird zudem der **Data Governance Act (DGA)** wirksam, der am 23. September 2022 in Kraft trat. Sein Ziel ist die Schaffung eines erleichterten Rahmens für die gemeinsame Nutzung von Daten. Ein europäisches Datenaustauschmodell soll zur Förderung der künstlichen Intelligenz einen Datenaustausch zwischen verschiedenen Branchen über Ländergrenzen hinweg ermöglichen. Bürger sollen ihre personenbezogenen Daten für bestimmte Zwecke spenden können. Zudem soll der Zugang zu Daten der öffentlichen Hand erleichtert werden. Datenvermittlungsdienste müssen in einem Register aufgeführt sein, damit interessierte Bürger sich von deren Vertrauenswürdigkeit überzeugen können.

Weiter voranschreiten dürfte 2023 auch die CSAM-Verordnung, die die EU-Kommission im Mai 2022 vorgelegt hat. **CSAM** steht für **Child Sexual Abuse Material**, also Kinderpornografie. Hosting- und Kommunikationsanbieter sollen danach Risikoeinschätzungen vornehmen und Maßnahmen zur Risikoreduzierung treffen. Sie werden dabei überwacht durch nationale Aufsichtsbehörden, denen besondere Befugnisse etwa in Bezug auf die Sicherstellung und Sperrung entsprechender Inhalte zustehen sollen.

Zu erwartende Neuerungen im IT-Recht 2023			
Kürzel	Name	in Kraft ab/seit	wirksam ab
AI Act	Artificial Intelligence Act	voraussichtlich 2023, spätestens 2024 (auch ein Scheitern ist nicht auszuschließen)	voraussichtlich nicht vor 2025, nach aktuellem Stand 24 Monate nach Inkrafttreten
CRA	Cyber Resilience Act	2023	24 Monate nach Inkrafttreten; einige erste Pflichten jedoch bereits 12 Monate nach Inkrafttreten
CSAM	„Child Sexual Abuse Material“-Verordnung	voraussichtlich 2023	voraussichtlich 6 Monate Umsetzungsfrist ab Inkrafttreten
DGA	Data Governance Act	23. September 2022	24. September 2024
DMA	Digital Markets Act	1. November 2022	2. Mai 2023
DORA	Digital Operational Resilience Act	verabschiedet am 10. November 2022; Inkrafttreten 20 Tage nach Veröffentlichung im EU-Amtsblatt	Jahreswechsel 2024/2025
DSA	Digital Services Act	16. November 2022	17. Februar 2024

Zu erwartende Neuerungen im IT-Recht 2023			
Kürzel	Name	in Kraft ab/seit	wirksam ab
	Ecodesign-Vorgaben, „Recht auf Reparatur“	2023	21 Monate Umsetzungsfrist ab Inkrafttreten
ECA	European Chips Act	voraussichtlich 2023	noch in Diskussion
ePVO	E-Privacy-Verordnung	eventuell 2023	nicht vor 2025
	EU-US Privacy Shield 2.0	eventuell 2023	
LksG	Lieferkettengesetz	1. Januar 2023	mit Inkrafttreten
MaRisk; BAIT	Mindestanforderungen an das Risikomanagement; Bankaufsichtsrechtliche Anforderungen an die IT	voraussichtlich 2023	
MiCA	Markets in Crypto-Assets	2023	18 Monate nach Inkrafttreten; voraussichtlich 2024
NIS2	Directive on Security of Network and Information Systems	voraussichtlich 2023, benötigt noch Zustimmung der EU-Staaten	voraussichtlich 2024, spätestens 2025

Abuse-Material: finden, löschen, berichten

Verfahren und Techniken zum Aufspüren kinderpornografischer Inhalte sollen bestimmten Vorgaben entsprechen, so datenschutzfreundlich und so wenig fehleranfällig wie möglich sein. Weitere Vorgaben soll ein noch zu schaffendes EU Centre

on Child Sexual Abuse (EU Centre) veröffentlichen. Zusätzlich gibt es für die verantwortlichen Unternehmen Berichtspflichten. Sie müssen entsprechende Inhalte löschen oder den Zugang zu ihnen effektiv unterbinden, wenn die Inhalte außerhalb der EU gehostet werden. Die Aufsichtsbehörden können Anordnungen treffen, denen unverzüglich Folge zu leisten ist.

App-Stores werden verpflichtet, den Download von Apps zu verhindern, die Kinder „einem hohen Risiko der Anwerbung [...] aussetzen können“. Das EU Centre steht dabei den Diensteanbietern, den einzelstaatlichen Ermittlungsbehörden sowie Europol, den EU-Mitgliedsstaaten und den Opfern beratend und unterstützend zur Seite. Wann die CSAM-Richtlinie verabschiedet werden wird, ist offen. Zuletzt hatten sich der Europäische Datenschutzbeauftragte und der Europäische Datenschutzausschuss kritisch geäußert. Sie werten die geplanten Regelungen als nicht vereinbar mit der Datenschutz-Grundverordnung und den freiheitlichen Grundrechten. Die emotionale Diskussion wird 2023 fortgesetzt werden.

Ab 1. Januar 2023 gilt das **Lieferkettengesetz**, zunächst für Unternehmen mit mehr als 3000 und ab 2024 auch für Unternehmen mit weniger als 1000 Beschäftigten. Es gilt zwar nicht ausschließlich für die IT-Branche, allerdings versprechen sich Marktbeobachter dort ein Umsatzwachstum, geht es doch um Automatisierung, Plattform as a Service, Supply-Chain-Management sowie Blockchain-Technologien. Ungeachtet der gesetzlichen Vorgaben dürfte die Diskussion um Diversifizierung der Beschaffung von Produkten, Rohstoffen und dergleichen auch 2023 anhalten.

Fazit

Aus IT-rechtlicher Sicht wird es das Jahr 2023 in sich haben. Die EU ist sehr umtriebig und wird zahlreiche Gesetzesvorhaben umsetzen. Auf Unternehmen aller Branchen kommen zahlreiche neue Vorgaben zu, etwa bei der Cybersicherheit. Einige der

Gesetzeswerke werden erst in den Jahren 2024 oder 2025 greifen. Zur Vorbereitung bleibt Unternehmen dennoch wenig Zeit. Denn ab Wirksamwerden der verschärften Vorgaben greifen signifikante Bußgelder nach dem Vorbild der Datenschutz-Grundverordnung. In manchen Fällen drohen auch Abmahnungen durch Verbände und Konkurrenten.

Ein Neujahrswunsch vieler betroffener Unternehmen für 2023 dürfte allerdings nicht in Erfüllung gehen: Es steht nicht zu erwarten, dass es vor der Europawahl 2024 noch zu einer Überarbeitung und Änderung der Datenschutz-Grundverordnung kommen wird. Hoffen darf man aber auf einen EU-US Privacy Shield 2.0 für die rechtssichere Übermittlung personenbezogener Daten in die USA. Hierzu wie auch in anderen Bereichen wird es auch im kommenden Jahr interessante und bedeutsame Gerichtsurteile geben, nicht zuletzt des Europäischen Gerichtshofs. Prosit 2023! (ur@ix.de)

1. Quellen

2. [Tobias Haar; EU will digitale Märkte regulieren; iX 9/2022, S. 80](#)
3. [Die im Text angesprochenen Gesetzesvorhaben sind über \[ix.de/zqe9\]\(https://www.ix.de/zqe9\) zu finden.](#)



Tobias Haar

ist Rechtsanwalt mit Schwerpunkt IT-Recht bei Vogel & Partner in Karlsruhe. Er hat zudem Rechtsinformatik studiert und hält einen MBA.

Kündigungsbutton: zahlreiche Abmahnungen

Kündigungsbutton: zahlreiche Abmahnungen

Seit Juli 2022 besteht eine Pflicht, Verbrauchern das Kündigen online abgeschlossener Verträge zu erleichtern. Die Verbraucherzentrale Bayern hat 840 Webseiten daraufhin untersucht und erhebliche Mängel festgestellt, die zu 154 Abmahnungen wegen Rechtsverstößen geführt haben. In einigen Fällen wird es zu Gerichtsverfahren kommen.

In zahlreichen Fällen war der vorgeschriebene Kündigungsbutton gar nicht vorhanden, in anderen Fällen war er versteckt und nicht wie gesetzlich gefordert „gut auffindbar“. Verstöße lagen auch gegen die Pflicht vor, eine gut lesbare Schaltfläche mit der Aufschrift „jetzt kündigen“ oder einer gleichwertigen Bezeichnung vorzuhalten. *Tobias Haar* (ur@ix.de)

Betrüger bestehlen sich gegenseitig

Sicherheitsexperten von Sophos

analysierten drei Untergrundforen und deren Schlichtungsräume für Streitigkeiten. Fazit: Wenn zwei Kriminelle sich streiten, freut sich die Verteidigung, die dadurch wertvolle Informationen erhält.



Markt + Trends | IT-Sicherheit

Dass die Schattenwelt der Internetkriminellen genauso arbeitsteilig agiert wie die „richtige“ Wirtschaft, ist seit einigen Jahren bekannt. Sicherheitsforscher von Sophos X-Ops veröffentlichen nun im ersten Teil einer vierteiligen Serie

neue Details (siehe ix.de/zey7). So verfügen die untersuchten Untergrundforen Exploit und XSS, zwei russischsprachige Cybercrime-Foren für Access as a Service (AaaS), und die englischsprachige, auf Datenlecks spezialisierte Plattform BreachForums mit Marktplatzfunktion über spezielle Schlichtungsräume zur Beilegung von Streitigkeiten. Dort können Nutzer Betrug, Angriffe und Abzocker melden.

Die „Betrüger betrügen Betrüger“-Masche scheint lukrativ zu sein: In einem Zeitraum von zwölf Monaten analysierte Sophos X-Ops rund 600 Betrugsfälle, bei denen die Bedrohungsakteure allein in diesen drei Foren mehr als 2,5 Millionen US-Dollar aneinander verloren.

Geld ist den Forschern von Sophos zufolge nicht die einzige Motivation, die die Kriminellen gegeneinander agieren lässt. Auch persönliche Streitigkeiten, Rivalitäten zwischen den Akteuren oder auch der Wunsch, den Ruf des anderen zu beschädigen oder den eigenen zu verbessern, gehören zu den Ursachen.

Scam Report Template:
by Baphomet - Saturday March 19, 2022 at 09:43 PM

March 19, 2022, 09:43 PM

Administrator

ADMINISTRATOR

Title of Thread: Scam Report against (username) | (amount)\$

Name: (Of the user who scammed you, also please give a link to his / her profile or just his UID)

Product: (Reference product name, link to thread, proof of purchase)

How did you get scammed: (Include details of the scam, screenshots, etc. The more information the better.)

Time of scam: (Tell us approximately when you got scammed)

It's important that both parties respond in a timely and respectful manner.

Sophos X-Ops

Im Untergrund wie im echten Leben: Beschwerde führen per Formular. *Sophos*

Die Angriffe gingen über das übliche „Abzocken und Verschwinden“ hinaus. Die Forscher sahen Empfehlungsbetrügereien, vorgetäuschte Datenabflüsse und gefälschte Tools, Phishing, URL-Hijacking, „alt rep“-Betrug (das Verfälschen von Reputationswerten durch Einsatz von „Sockpuppets“, also Fake-Accounts), falsche Bürgen, Erpressung, nachgemachte Konten und Backdoors. Auch konnten

die Forscher Fälle beobachten, in denen sich betrogene Bedrohungsakteure wiederum an ihren Betrügern rächen.

Die Sicherheitsforscher fanden überdies Indizien für langfristigen, groß angelegten Betrug in Form von neunzehn Websites, alle von derselben Person oder Gruppierung erstellt, die kriminellen Marktplätzen täuschend ähnlich sehen. Sie fordern von neuen Nutzern eine Aktivierungsgebühr in Höhe von 100 Dollar.

Verteidiger

Theoretisch könnte es der Allgemeinheit völlig gleichgültig sein, wie Kriminelle und Betrüger zueinander stehen oder miteinander umgehen. Aber, erläutert Matt Wixey, Senior Threat Researcher bei Sophos, „da Kriminelle oft viele Beweise vorlegen müssen, wenn sie über die Betrügereien berichten, denen sie selbst zum Opfer gefallen sind, liefern sie eine Fülle von taktischen und strategischen Informationen über ihre Operationen – eine bisher ungenutzte Ressource“. Diese Schlichtungsberichte vermitteln außerdem einen Einblick in die Prioritäten der Angreifer, ihre Rivalitäten und Allianzen, „und, ironischerweise, wie anfällig sie für die gleichen Arten von Täuschung sind, die sie gegen ihre Opfer einsetzen“, so Wixey. (ur@ix.de)

ix.de/zey7

- [The scammers who scam scammers on cybercrime forums: Part 1](#)
- [Folien des Black-Hat-Vortrags von Sophos](#)
- [BMI-Papier: Strategie zur Bekämpfung der Schweren und Organisierten Kriminalität](#)
- [NSA-Empfehlungen für Entwickler zum Absichern der Supply Chain](#)
- [Projekt Sigstore – Software Signing for Everybody](#)

- [Konzept von Sigstore](#)
- [verinice.veo DSMS](#)
- [Playlist der Vorträge der Black Hat 2022](#)
- [Aagon Bitlocker-Management](#)



The scammers who scam scammers on cybercrime forums: Part 1

A shadowy sub-economy is more than just a curiosity – it's booming business, and also an opportunity for defenders. In the first of a four-part series, we look at the forums involved, and how they ...

Die Betrüger, die Betrüger in Cybercrime-Foren betrügen: Teil 1

Eine Schattenwirtschaft ist mehr als nur eine Kuriosität – sie ist ein boomendes Geschäft und auch eine Chance für Verteidiger. Im ersten einer vierteiligen Serie betrachten wir die beteiligten Foren und wie sie mit Betrügern umgehen, die Betrüger betrügen

Geschrieben von [Matt Wixey](#)

[07. Dezember 2022](#)

[Bedrohungsforschung](#) [AaaS](#) [BreachForums](#) [Exploit](#) [RaidForums](#)
[Marktplätze](#) [empfohlene](#) [Betrug](#) [Sophos](#) [X-Ops](#) [XSS](#)

Auf kriminellen Marktplätzen lauert an jeder Ecke ein Betrug. Bereits 2009 [wies Microsoft darauf hin, dass die Untergrundwirtschaft voller Unehrlichkeit](#) sei, und 2017 berichtete Digital Shadows über eine Datenbank von „Rippern“

(Betrüger, die Kriminelle betrügen), die von Marktplatzbenutzern erstellt wurde. In [unserer jüngsten Berichterstattung über Genesis Market](#) haben wir mindestens eine betrügerische Imitation von Genesis festgestellt, die darauf abzielt, naive Möchtegern-Cyberkriminelle (und möglicherweise unerfahrene Sicherheitsforscher und Journalisten) von ihrem Geld zu trennen.

Aber im Allgemeinen hat das Thema nicht viel Aufmerksamkeit erhalten. Warum sollte es denn auch? Wenn Betrüger Kriminelle ins Visier nehmen, umso besser, oder? Zumindest greifen sie sich gegenseitig an, nicht Organisationen oder die breite Öffentlichkeit.

Wir dachten, dass da noch mehr dahintersteckt, also verbrachten wir ein paar Wochen damit, Betrüger zu untersuchen, die Betrüger in drei prominenten Cybercrime-Foren betrügen – eine Recherche, die unserer Meinung nach noch nie zuvor durchgeführt wurde. Und wir fanden fünf überraschende Dinge.

1. Es ist ein großes Geschäft – eine Subökonomie für sich. In den letzten 12 Monaten haben Cyberkriminelle allein in diesen drei Foren über 2,5 Millionen US-Dollar durch Betrug verloren. Tatsächlich ist es ein so lange bestehendes und prominentes Problem, dass Forenadministratoren spezielle „Schlichtungsräume“ eingerichtet haben, in denen Benutzer Betrug, Angriffe und Ripper melden können.

2. Geld ist nicht das einzige Motiv, und es sind nicht nur niederrangige Bedrohungsakteure beteiligt. Persönliche Probleme, Rivalitäten und der Wunsch, den Ruf zu zerstören (oder manchmal zu verbessern), können alle zu Betrug führen. Und es sind nicht nur kleine Gauner. Wir sahen prominente Bedrohungsakteure, die entweder des Betrugs beschuldigt wurden oder selbst Opfer von Betrug wurden.

3. Die Angriffe gehen über das übliche „Rip-and-Run“ hinaus.

Wir sahen Verweis-Nachteile, gefälschte Datenlecks und Tools, Typosquatting, Phishing, „Alt-Rep“-Betrug (die Verwendung von Sockenpuppen, um die Reputationswerte künstlich aufzublähen), gefälschte Bürgen, Erpressung, imitierte Konten und Backdoor-Malware. Wir haben sogar Fälle gefunden, in denen sich Bedrohungsakteure rächen, indem sie die Betrüger betrügen, die sie betrogen haben.

4. Wir haben Beispiele für langfristigen, groß angelegten Betrug gefunden. Eine der größten Überraschungen kam, als wir uns mit dieser nachgeahmten Genesis-Seite befassten. Mit einiger Detektivarbeit entdeckten wir neunzehn weitere Websites, die alle von derselben Person oder Gruppe erstellt wurden, alle kriminelle Marktplätze imitierten und alle darauf abzielten, Benutzer dazu zu verleiten, eine „Aktivierungsgebühr“ von über 100 US-Dollar zu zahlen. Wir wissen nicht genau, wer hinter all diesen Seiten steckt, aber wir haben versuchsweise Links zu einem Drogenhändler entdeckt, der auf mehreren dunklen Websites operiert.

So weit, so *Schadenfreude* – aber die große Frage ist immer noch: wen interessiert das? Warum spielt es eine Rolle, wenn sich Kriminelle gegenseitig angreifen? Hier wird es wirklich faszinierend.

5. Betrugsberichte sind eine reichhaltige und wenig erforschte Informationsquelle. Bedrohungsakteure sind sich bewusst, dass kriminelle Foren überwacht werden, und setzen daher häufig auf gute Betriebssicherheit. Wenn sie selbst Opfer von Verbrechen sind – nun ja, nicht so sehr. Da Forenregeln Beweise für Betrugsvorwürfe verlangen, posten Angreifer, denen Unrecht getan wurde, oft gerne Screenshots von privaten Gesprächen und Quellcode, Identifikatoren, Transaktionen, Chatprotokollen und detaillierte Berichte über Verhandlungen, Verkäufe und Fehlerbehebung.

Diese versteckte Subwirtschaft ist nicht nur eine Kuriosität. Es gibt uns Einblicke in die Forumskultur; wie

Bedrohungsakteure kaufen und verkaufen; ihre taktischen und strategischen Prioritäten; ihre Rivalen und Allianzen; ihre Anfälligkeit für Täuschung – und spezifische, diskrete Informationen über sie.

In den nächsten Wochen werden wir die Ergebnisse unserer ausführlichen Untersuchung zu diesem Thema teilen – beginnend mit einem Überblick über die beteiligten Foren, wie sie mit Betrug umgehen, wer wen betrügt und die Größe der Subwirtschaft.

Sie können sich auch [unseren Black-Hat-Vortrag](#) zu dieser Forschung ansehen.

Willkommen im Dschungel

Um unsere Untersuchung einzuleiten, haben wir Betrügereien in zwei der ältesten und bekanntesten russischsprachigen Cybercrime-Foren, Exploit und XSS, untersucht. Wir haben auch Betrügereien von BreachForums, dem Nachfolger von RaidForums, das im April 2022 gestartet wurde, aufgenommen.

Die Foren

Exploit ist relativ exklusiv und ein beliebter Marktplatz für [Access-as-a-Service \(AaaS\)-Angebote](#), bei denen [Initial Access Brokers \(IABs\) den Zugang zu kompromittierten Netzwerken verkaufen](#). Aber Bedrohungsakteure kaufen und verkaufen dort auch viele andere illegale Inhalte – Malware, Datenlecks, Infostealer-Protokolle, Anmeldeinformationen und mehr. In der Vergangenheit besuchten Ransomware-Gruppen und -Partner Exploit, obwohl dies nach dem Angriff auf die Colonial Pipeline im Jahr 2021 verdeckter wurde, als [sowohl Exploit als auch XSS Ransomware-Diskussionen öffentlich untersagten, um negative Aufmerksamkeit zu vermeiden](#). Heutzutage wird die Rekrutierung von Ransomware-Affiliates in beiden Foren fortgesetzt, obwohl dies eher unter dem Deckmantel von Euphemismen wie „Pentester“ erfolgt.

XSS, früher bekannt als DaMaGeLaBs, ist ebenfalls gut etabliert, obwohl die Mitgliedschaft weniger exklusiv ist als Exploit. Es hostet auch viele AaaS-Angebote und verschiedene andere Inhalte.

Schließlich ist BreachForums der Nachfolger von RaidForums, einem Marktplatz, der sieben Jahre lang lief, [bevor er Anfang 2022 von den Strafverfolgungsbehörden beschlagnahmt wurde](#). BreachForums ist wie RaidForums ein englischsprachiges Cybercrime-Forum und ein Marktplatz, der sich auf Datenlecks spezialisiert hat, darunter personenbezogene Daten, Kreditkarten, Anmeldeinformationen und Ausweisdokumente.

Alle drei Seiten haben dedizierte Schlichtungsräume – Exploit (mit ungefähr 2500 gemeldeten Betrügereien) und XSS (mit ungefähr 760) haben sie seit Mitte der 2000er Jahre und BreachForums seit ihrer Gründung im April 2022. Andere kriminelle Marktplätze, wie Verified, haben sie Sie auch.

Tatsächlich hat Exploit zwei Räume – einen für offene Ansprüche und einen anderen, der als „Schwarze Liste“ bezeichnet wird und bestätigte Betrugsfälle dokumentiert.



Abbildung 1: Arbitration-Bereich von Exploit

Zusätzlich zu einem speziellen Schlichtungsraum führt XSS auch eine lange „Ripper-Liste“, einen Index von Betrugsseiten.



Abbildung 2: Die Ripper-Liste von XSS

Eine Übersicht über Betrugsstatistiken

Wir haben uns alle Betrugsberichte der letzten 12 Monate angesehen, in denen Geldbeträge angegeben wurden. (Mit BreachForums gingen wir zurück zum ersten aufgezeichneten Betrug, da das Forum noch nicht so lange existiert.)

	Exploit (offene Ansprüche)	Exploit („Schwarze Liste“)	XSS	Verletzungsforen
Ansprüche	211	236	120	21
Gesamtmenge	\$1,021,998	\$863,324	\$509,901	\$143,722
Bedeuten	\$4,843.54	\$3,658	\$4,249.18	\$6,843.90
Modus	\$1000	\$500	\$150	\$500
Median	\$600	\$500	\$500	\$200
Bereich	\$15 – \$160,000	\$5 – \$150,000	\$10 – \$160,000	\$2 – \$134,000

Tabelle 1: Eine Zusammenfassung von 12 Monaten Betrugsmeldungen (alle Beträge in USD)

Dies ist zwar nur eine Momentaufnahme, gibt uns aber einige nützliche Einblicke. Erstens beträgt der durch Betrug verlorene Gesamtbetrag (und denken Sie daran, dass dies nur Betrugsberichte betrifft, in denen bestimmte Beträge erwähnt werden – manche tun dies nicht) 2.538.945 \$. Das ist eine beträchtliche Menge, wenn man bedenkt, dass es sich nur um drei Foren handelt.

Zweitens ist Exploit das Schlimmste für Betrug, sowohl in Bezug auf die Anzahl der Berichte als auch auf das Geld, das Betrügern verloren geht. Es hat etwa doppelt so viele Mitglieder wie XSS und kann aufgrund seines guten Rufs auch mehr Betrüger anziehen.

Drittens ist der durchschnittliche als gestohlen gemeldete Betrag in allen drei Foren ähnlich, ebenso wie die Bandbreite – was darauf hindeutet, dass das Ausmaß der Betrügereien unabhängig vom Forum gleich ist.

Opfer haben Betrugsmeldungen für nur 2 US-Dollar eingereicht; Angreifer scheinen genauso empört über den Diebstahl ihres Geldes zu sein wie alle anderen, egal wie hoch der Betrag ist.

Am oberen Ende gehen die Betrugereien auf allen drei Marktplätzen in den sechsstelligen Bereich, obwohl dies die Ausnahmen sind. Viele Betrugereien bringen relativ unbedeutende Beträge ein.



Abbildung 3: Niedrige Schadenssummen im XSS-Schlichtungsraum



Abbildung 4: Niedrige Forderungsbeträge im Schlichtungsraum von BreachForums



Abbildung 5: Ein Beispiel für einen größeren Betrugsanspruch auf Exploit (130.000 \$). Beachten Sie die vielen Details in diesem Betrugsfall, der Informationen über Verhandlungen und Projekte enthält

Bevor wir uns mit dem Schlichtungsverfahren befassen, lohnt es sich zu untersuchen, warum Betrug so weit verbreitet ist. Bereits 2009 argumentierte Microsoft, dass die illegale Cyberkriminalität keine „kriminelle Utopie des leichten Geldes“ sei, sondern ein „Zitronenmarkt“, auf dem die Anwesenheit von Rippnern effektiv eine Steuer auf jede Transaktion einführte.

Auch wenn sich die Zeiten geändert haben und Cyberkriminalität immer mehr zur Ware geworden ist, sind kriminelle Marktplätze immer noch der perfekte Nährboden für Betrüger und Ripper. Es gibt keinen Rückgriff auf die Strafverfolgung; es ist eine (halb) anonyme Kultur, die Privatsphäre betont; Websites sind so exklusiv, dass zumindest ein gewisses Maß an implizitem Vertrauen besteht; Sie werden von Kriminellen bevölkert, die sich wohl kaum als potenzielle Opfer betrachten und daher möglicherweise weniger auf der Hut vor Betrug sind. es ist ein offener Markt ohne Regulierung oder Qualitätssicherung; Transaktionen werden mit Kryptowährungen durchgeführt, die effektiv unauffindbar gemacht werden können; und

Sicherheitsvorkehrungen wie Bürgen sind optional (und können, wie wir im nächsten Teil unserer Serie sehen werden, selbst in den Dienst von Betrugereien gestellt werden).

Was unternehmen kriminelle Marktplätze gegen Betrug?

Die Administratoren krimineller Foren sind sich bewusst, dass Betrug ein Problem darstellt. Zusätzlich zu den Schlichtungsstellen verfügen die meisten Marktplätze über sichtbare Warnungen vor Betrügern und befürworten die Verwendung von Bürgen (manchmal auch als „Zwischenhändler“ oder „Mittelsmänner“ bezeichnet) während des Verkaufs – eine Form der Treuhand.



Abbildung 6: Eine Warnung vor Betrug auf der Titelseite von BreachForums

Andere Foren gehen weiter. Verified zum Beispiel warnt Benutzer ausdrücklich vor gefälschten Links zu seinem Forum und befürwortet die Verwendung eines benutzerdefinierten Plugins, um solche Betrugereien zu erkennen:



Abbildung 7: Betrugswarnung von Verified

In ähnlicher Weise veröffentlicht BreachForums eine Liste aller seiner legitimen Domänen sowie einen monatlichen „Transparenzbericht“, um zu bestätigen, dass die Website und die zugehörige Infrastruktur unter seiner Kontrolle bleiben und nicht kompromittiert wurden (obwohl dies wahrscheinlich auch eine Vorsichtsmaßnahme ist [Maßnahme aufgrund dessen, was mit RaidForums passiert ist](#)):



Abbildung 8: Einzelheiten zum monatlichen Transparenzbericht von BreachForums

Aber Schlichtungsstellen sind die Hauptmethode für den Umgang mit Betrug. Der Prozess ist relativ einfach. Benutzer, die einen Betrug melden möchten, müssen einen neuen Thread erstellen, den Benutzer anrufen, der sie angeblich betrogen hat, und so viele Details wie möglich über den Vorfall angeben. BreachForums stellt hierfür eine Vorlage bereit, während XSS lediglich die erforderlichen Details auflistet.



Abbildung 9: Vorlage für Betrugsberichte von BreachForums



Abbildung 10: Die in XSS-Betrugsberichten erforderlichen Daten: Benutzername, Link zum Profil, Kontaktdaten, Beweise (Chatprotokolle, Screenshots, Brieftaschen, Überweisungen), alle zusätzlichen Informationen

Ein Moderator überprüft dann den Bericht, bittet um weitere Informationen, falls erforderlich, markiert den Angeklagten und gibt ihm eine Frist für die Antwort (normalerweise 24 Stunden, kann aber zwischen 12 und 72 Stunden liegen).



Abbildung 11: Ein Exploit-Moderator gibt einem beschuldigten Betrüger 24 Stunden Zeit, um auf einen Vorwurf zu reagieren

Der Angeklagte kann die Forderung akzeptieren, in diesem Fall leistet er dem Opfer Wiedergutmachung. Das ist selten. Häufiger bestreitet der Angeklagte die Behauptung (in diesem Fall entscheidet der Moderator) oder antwortet überhaupt nicht (in diesem Fall kann er vorübergehend oder dauerhaft aus dem Forum ausgeschlossen werden).



Abbildung 12: Ein umstrittener Anspruch auf XSS in Bezug auf AaaS-Angebote

Bei strittigen Behauptungen kann der Moderator für eine Partei entscheiden oder entscheiden, dass aufgrund fehlender Beweise

kein Fall zu beantworten ist. In einigen Fällen erhalten eine oder beide Parteien Verwarnungen oder vorübergehende oder dauerhafte Sperren.



Abbildung 13: Der Administrator von BreachForums schließt einen Betrugsbericht aufgrund fehlender Beweise



Abbildung 14: Ein umstrittener Anspruch auf Exploit bezüglich eines Crypters zur Verwendung mit [Remcos](#)

Diese Diskussionen sind manchmal zivil und werden gütlich zur Zufriedenheit beider Parteien beigelegt. Wir haben ein Beispiel notiert, bei dem der Schiedsrichter entschied, dass der Angeklagte 50 % des geforderten Betrags zurückzahlen sollte:



Abbildung 15: Ein Exploit-Moderator gibt dem Angeklagten 24 Stunden Zeit, um 50 % des geforderten Betrags zurückzuzahlen

In einem Fall entschädigte der Administrator von BreachForums sogar ein Betrugsoffer aus eigener Tasche:



Abbildung 16: Der Administrator von BreachForums entschädigt ein Betrugsoffer persönlich mit 200 US-Dollar

Betrugsberichte enden jedoch häufiger in Beleidigungen und Gegenanschuldigungen. In einigen Fällen wurden die mutmaßlichen Opfer später selbst wegen Betrugs gesperrt.



Abbildung 17: Ein Betrugsbericht über Exploit führt dazu, dass der Ankläger den Ankläger des Betrugs beschuldigt

Folgen

Verbote (und in geringerem Maße Verwarnungen) scheinen das häufigste Ergebnis in Schiedsverfahren zu sein, aber BreachForums verfolgt einen etwas anderen Ansatz. Vielleicht, um zukünftige Betrüger abzuschrecken, veröffentlichen die Moderatoren die Registrierungs-E-Mail-Adressen und Registrierungs- und zuletzt gesehenen IP-Adressen gesperrter Benutzer, wodurch sie teilweise doxiert werden:



Abbildung 18: Ein Beispiel eines gesperrten Benutzers, komplett mit veröffentlichter Registrierungs-E-Mail-Adresse, Registrierung und letzten bekannten IP-Adressen

Wir haben einige Fälle von Serienbetrügern bemerkt, die nach einer Sperrung einfach ein neues Profil mit einer neuen Identität erstellten, eine neue Registrierungsgebühr zahlten und wieder mit dem Betrügen begannen.

Nicht nur kleine Gauner

Wir haben einige Beispiele notiert, an denen prominentere Bedrohungsakteure beteiligt waren. Hier ist zum Beispiel ein merkwürdiger Fall, der nicht so sehr ein Betrug war, sondern einen Benutzer betraf, der im Namen eines Opfers mit der Conti-Ransomware-Gruppe verhandeln wollte:



Abbildung 19: Ein Benutzer erhebt eine Schiedsklage, um zu versuchen, mit der Conti-Gruppe über die Entschlüsselung der Vermögenswerte eines Unternehmens zu verhandeln

Dieser Bericht wurde von Exploit-Moderatoren geschlossen, da er sich auf Ransomware bezog, die angeblich in diesem Forum verboten ist. Interessant ist jedoch, dass der Beschwerdeführer selbst ein Bedrohungsakteur zu sein scheint und dem Exploit-Forum über drei Jahre lang beigetreten war, bevor er den oben genannten Anspruch geltend machte – mit

mehreren Beiträgen, in denen er sein Interesse am Kauf von Daten bekundete. Ihre Beziehung zu Contis Opfer in diesem Fall ist nicht klar.



Abbildung 20: Einige der früheren Beiträge des Beschwerdeführers im Exploit-Forum

Ein weiterer Fall betraf „Alan Wake“ (ein Name aus einem Videospiel), der den letzten [Wettbewerb auf XSS](#) gesponsert hatte und zuvor [von einem Lockbit-Betreiber beschuldigt wurde, der Anführer der Ransomware-Gruppen Conti und BlackBasta zu sein](#). Ein Benutzer beschuldigte Alan Wake, sein Gehalt nicht gezahlt zu haben, weil er „Verkehr aus Muscheln gemacht“ habe:



Abbildung 21: Der XSS-Betrugsbericht gegen „Alan Wake“

Alan Wake bestritt den Vorwurf, und der Fall wurde vom Administrator geschlossen und der Beschwerdeführer gesperrt – nicht wegen Betrugs, sondern wegen „Beleidigungen, Angriffen, Drohungen usw.“ und „äußerst unangemessenem Verhalten“.

Schließlich wurde All World Cards (ebenfalls ein früherer Sponsor von XSS-Wettbewerben), eine prominente Carding-Gruppe, selbst Opfer eines Betrugs mit einer gefälschten Schwachstelle und verlor 2000 USD.



Abbildung 22: Die Gruppe All World Cards meldet einen Betrug, bei dem sie 2000 Dollar verloren hat

Wenn es eine Erkenntnis aus all dem gibt, dann die, dass kein Benutzer immun ist; Jeder Handel in kriminellen Foren birgt ein inhärentes Betrugsrisiko. Obwohl es sowohl proaktive (Warnungen, Plugins, Garanten) als auch reaktive (Schlichtungsstellen) Maßnahmen gibt, sind Betrüger nicht nur üblich, sondern – nach den von uns gesammelten Daten zu urteilen – oft erfolgreich. Einer der Gründe für ihren Erfolg

ist die schiere Vielfalt der Betrügereien, die sie ziehen.

Im zweiten Teil unserer Untersuchung, der nächste Woche um diese Zeit (Mittwoch, 14. Dezember) erscheinen wird, behandeln wir die verschiedenen Arten von Betrug, die wir beobachtet haben.



Managed Security Services: 4 kritische Fragen, die Sie stellen sollten

Die Landschaft der Cybersicherheitsbedrohungen ist unglaublich volatil. Cyberkriminelle gehen immer professioneller vor, spezialisieren sich zunehmend und treten sogar in Konkurrenz zu anderen Grup...

Managed Security Services: 4 kritische Fragen, die Sie stellen sollten

Verfasst von [Jörg Schindler](#)

[24. November 2022](#)

Die Landschaft der Cybersicherheitsbedrohungen ist unglaublich volatil. Cyberkriminelle gehen immer professioneller vor, spezialisieren sich zunehmend und treten sogar in Konkurrenz zu anderen Gruppierungen. In der Folge sind Unternehmen innerhalb von Monaten, Wochen oder Tagen – manchmal sogar gleichzeitig – nicht nur einmal sondern immer wieder Angriffen ausgesetzt.

Der weltweite Arbeitskräftemangel im Bereich Cybersicherheit verschärft diese Herausforderungen. Weltweit hat sich die Personallücke im Bereich Cybersicherheit im Jahr 2022 nach

Angaben der „2022 Cybersecurity Workforce Study by (ISC)²“ um 26,2 % erhöht, mit insgesamt mehr als drei Millionen offenen Stellen. Während einige Regionen besser abschneiden als andere – wie beispielsweise Lateinamerika, das die Lücke um 26,4 % schloss – bergen die verbleibenden Engpässe immer noch nationale Sicherheitsrisiken.

Cyberkriminelle sind immer aktiv, und Sicherheitsteams müssen es auch sein. Viele Organisationen, die nicht über die erforderlichen Ressourcen verfügen, um selbst immer komplexere Cyberbedrohungen zu erkennen und darauf zu reagieren, entscheiden sich für die Nutzung von Cybersecurity-as-a-Service (CSaaS), um proaktive Abwehrmaßnahmen zu implementieren. Beim CSaaS-Modell setzen Unternehmen externe Spezialisten ein, um kritische Cybersicherheitsanforderungen zu erfüllen, wie z. B. Bedrohungsüberwachung rund um die Uhr. Durch die Auslagerung oder Erweiterung von IT-Teams mit Managed Detection and Response (MDR)-Services als zentrales CSaaS-Angebot können Unternehmen dazu beitragen, Angriffe abzuschwächen, bevor sie auftreten. Jedes Unternehmen, das erwägt, den Sicherheitsbetrieb auszulagern, sollte Sicherheitsdienstpartnern diese vier Fragen stellen:

1. Welche Erfahrung haben sie in der Zusammenarbeit mit anderen Unternehmen in unserer Branche und Region?

Wenn der Anbieter mit anderen Organisationen in ihrer Branche und Region zusammenarbeitet, sollten diese über Erfahrungen aus erster Hand bei der Verteidigung gegen die spezifischen Bedrohungen verfügen, denen sie ausgesetzt sind.

2. Können sie unsere bestehenden Technologien verwalten und unterstützen?

Fragen sie, ob der CaaS-Anbieter auf ihren vorhandenen Sicherheitstechnologien aufbauen kann, oder ob sie das, was sie bereits im Einsatz haben, entfernen und ersetzen müssen. Der ideale Anbieter sollte in der Lage sein, mit den vorhandenen Technologielösungen zu arbeiten.

3. Wie ausgereift ist ihr Verständnis von neu auftretenden Cyberbedrohungen?

Kriminelle entwickeln sich häufig in den Taktiken, Techniken und Verfahren (TTPs) weiter, die sie verwenden, um Angriffe möglichst unbemerkt durchzuführen. Unternehmen sollten sehr sorgfältig darauf achten, dass ein potenzieller Anbieter die entsprechenden Ressourcen vorweisen kann, mit denen er eine qualitativ hochwertige Bedrohungsanalyse sowie schnelle Reaktion gewährleisten kann.

4. Kann die Lösung eines potenziellen Partners mit unserem Unternehmen skalieren und sich mit unseren Anforderungen weiterentwickeln?

Es ist von entscheidender Bedeutung, dass jeder potenzielle Partner in der Lage ist, den individuell wachsenden und sich entwickelnden Anforderungen gerecht zu werden und die Unternehmenssicherheit zusammen mit sich ändernden Anforderungen effektiv zu optimieren.

Mit einem starken CSaaS-Anbieter sind Unternehmen in der Lage, eine vollständig etablierte Sicherheitsstruktur mit proaktiven Abwehrmaßnahmen und [24/7-Unterstützung](#) zu realisieren. Dies gibt Unternehmen die Möglichkeit, ihre IT-Operationen kontinuierlich zu verbessern und Organisationsmodelle zu verfeinern, wodurch sie in einer äußerst volatilen Bedrohungslandschaft nicht nur überleben, sondern auch wachsen können.



Black Hat

Black Hat

Betrüger, die Betrüger betrügen,

Hacker, die Hacker hacken: Erkundung einer verborgenen Subökonomie in Foren und Marktplätzen für Cyberkriminalität

[Matt Wixey](#) | Leitender technischer Redakteur, Sophos
[Angela Gunn](#) | Senior Threat Researcher / Cybersecurity
Writer, Sophos

Datum : Mittwoch, 7. Dezember | 15:20-16:00 Uhr (Capital Suite
Zimmer 7/12 (Ebene 3))

Format : 40-Minuten-Briefings

Spuren : Menschliche Faktoren, Verteidigung Es ist kein Geheimnis, dass kriminelle Foren und Marktplätze mit schändlichen Aktivitäten vollgestopft sind. Aber hinter all den Initial Access Brokern, gestohlenen Daten und Malware gibt es eine versteckte, blühende Unterkategorie der Kriminalität, die unbemerkt bleibt: Bedrohungsakteure, die es auf andere Bedrohungsakteure abgesehen haben. Diese kannibalischen Kriminellen (wir nennen sie „Metaparasiten“: ein Parasit, dessen Wirt auch ein Parasit ist) sind ein so hartnäckiges und teures Problem, dass es spezielle Forenräume gibt – die Tausende von Posts enthalten und Jahre zurückreichen –, die dafür bestimmt sind, sie auf die schwarze Liste zu setzen und Betrug zu schlichten Beschwerden zwischen Benutzern und das Melden von nachgeahmten „Ripper“-Sites. In diesem Vortrag präsentieren wir eine neuartige Untersuchung über Betrüger, die Betrüger betrügen, und Hacker, die Hacker hacken, auf drei der etabliertesten und bekanntesten kriminellen Marktplätze. Wir untersuchen die Größe dieses schattigen Multi-Millionen-Dollar-Ökosystems; die Beweggründe von Metaparasiten; wie Schiedsverfahren funktionieren; und welchen Einfluss Metaparasiten auf die Kultur und den Betrieb der Marktplätze haben, auf denen sie tätig sind. Anschließend tauchen wir tief

in Fallstudien ein und betrachten die Techniken, die Metaparasiten verwenden, von altmodischem „Rip and Run“-Betrug und gefälschten Datenlecks bis hin zu ausgeklügelten Phishing-Kampagnen, Verweissbetrug, Typosquatting und Backdoor-Malware. Unterwegs decken wir einen groß angelegten, koordinierten und lukrativen Betrug auf, an dem ein Netzwerk von 15 gefälschten Marktplätzen beteiligt ist, und Fälle, in denen sich die Bedrohungsakteure rächen und die Betrüger betrügen, die sie betrogen haben. Sie könnten fragen: Wen kümmert es, wenn Kriminelle sich gegenseitig abzocken? Aber Metaparasiten bieten Analysten unbeabsichtigt einen Informationssegen, der es uns ermöglicht, beispiellose Einblicke in Verkäufe, Operationen, Verhandlungen und Identifikatoren zu gewinnen, die sonst verborgen bleiben würden – sowie in die Marktkultur, unterschiedliche Ebenen der Betriebssicherheit und Anfälligkeit für Täuschung und Sozialtechnik. Unser Vortrag wird auch dazu beitragen, Analysten und allgemein Neugierige davor zu schützen, versehentlich auf einige dieser Betrügereien hereinzufallen, wenn sie kriminelle Marktplätze untersuchen.

Präsentationsmaterial

- [Folien hier herunterladen](#)

EU verabschiedet NIS2-Richtlinie – Umsetzung bis 2024

Nach dem EU-Parlament hat auch der EU-Rat der Richtlinie über Maßnahmen zur Gewährleistung eines hohen gemeinsamen Sicherheitsniveaus von Netz- und Informationssystemen in der Union (NIS2) zugestimmt. Die Mitgliedsstaaten müssen sie bis Herbst 2024 in nationales Recht umsetzen. Sie soll die Resilienz und die Kapazitäten zur Reaktion auf Sicherheitsvorfälle sowohl des öffentlichen als auch des privaten Sektors und der EU als Ganzes weiter verbessern. Die NIS2-Richtlinie steht im Zusammenhang mit zahlreichen gesetzgeberischen Maßnahmen im Bereich IT- und

Cybersicherheit.

Ähnlich den Vorgaben der KRITIS-Verordnung werden Unternehmen betroffener Branchen verpflichtet, Risikomanagementmaßnahmen zu ergreifen und Meldepflichten zu beachten. Zu den Branchen zählen unter anderem Energie, Verkehr, Gesundheit und digitale Infrastruktur. Um sicherzustellen, dass nur mittlere und große Unternehmen von den Vorgaben erfasst werden, sieht die Richtlinie Schwellenwerte für ihre Anwendbarkeit vor.

Ziel der Richtlinie ist eine Harmonisierung der einschlägigen Bestimmungen in den einzelnen EU-Staaten durch Mindestvorgaben und Regeln zur wirksameren Zusammenarbeit zwischen den nationalen Behörden. Unter anderem bei Aspekten der Zusammenarbeit und Kooperation sowie den Anforderungen an das Cybersecurity-Risikomanagement geht die Richtlinie deutlich über die NIS1-Richtlinie hinaus. *Tobias Haar* (ur@ix.de)

Kurz notiert

Aagon veröffentlicht ein Produkt zum **BitLocker-Management**. Es bietet die zentrale Verwaltung des Windows-Bordmittels zur Festplattenverschlüsselung sowie Monitoring- und Reportfunktionen.

Seit Kurzem stehen 102 Vorträge der diesjährigen **Sicherheitskonferenz Black Hat** auf dem YouTube-Channel des Veranstalters zum Nachschauen bereit (siehe ix.de/zey7).

Der **verinice.veo-Datenschutzmanager** steht Interessierten in einer Einzelplatz-Betatestversion zur Verfügung. Mit dem Datenschutzmanagementsystem lassen sich die Vorgaben der DSGVO verwalten und ihre Umsetzung gewährleisten.

So funktioniert der TikTok-Algorithmus: Alles, was Sie wissen müssen

Wir lesen die Richtlinien für Ersteller und die Newsroom-Dokumentation von TikTok, damit Sie es nicht tun müssen. Erfahren Sie alles über TikTok und wie Marken es nutzen können.

Seit seiner Erstveröffentlichung im Jahr 2016 hat sich TikTok zur am [schnellsten wachsenden Social-Media-Plattform entwickelt](#) .

Wenn Sie mehr über die beliebte Kurzform-Video-App erfahren möchten und wie Ihre Marke sie nutzen kann, sind Sie hier genau richtig.

-Dokumentation von TikTok [Wir lesen die Richtlinien für Ersteller](#) und [die Newsroom](#) , damit Sie es nicht tun müssen.

In diesem Artikel teilen **wir alles** , was wir gelernt haben – zusammen mit Erkenntnissen von TikTokern, die sich mit der Steigerung der Aufrufe, der Einbindung von Benutzern und dem Aufbau einer Fangemeinde auskennen.

Finden Sie heraus, wie der TikTok-Algorithmus funktioniert, was Sie tun müssen, um ein erfolgreiches Video zu erstellen, und entscheiden Sie, ob TikTok Ihrer Marke zugute kommen kann.

In diesem Artikel:

- [Was ist der TikTok-Algorithmus?](#)
- [Was ist der For You-Feed?](#)
- [Wie funktioniert der TikTok-Algorithmus?](#)
- [Keywords und der TikTok-Algorithmus](#)
- [TikTok-SEO und Google](#)
- [TikTok-Algorithmus-Mythen entlarvt](#)
- [Was der Algorithmus nicht zeigt](#)
- [7 Tipps, um mit dem TikTok-Algorithmus zu arbeiten und viral zu werden](#)
- [Warum sich Marken für TikTok interessieren sollten](#)

Was ist der TikTok-Algorithmus?

Der TikTok-Algorithmus ist ein [Empfehlungssystem](#), das die Videos bestimmt, die Sie in Ihrer App sehen. Wie jeder gute Algorithmus funktioniert es, um relevante Inhalte zu bringen, die Ihnen gefallen, basierend auf Ihren Interessen.

Empfehlungssysteme werden überall in der digitalen Welt verwendet und sind nicht neu. Viele Plattformen, einschließlich Netflix, verwenden sie, um Ihnen die Inhalte bereitzustellen, die Ihnen am besten gefallen.

Der Zweck eines Empfehlungssystems besteht darin, Inhalte zu teilen, die Benutzer basierend auf ihren Vorlieben und denen von Personen mit ähnlichen demografischen Merkmalen mögen.

The TikTok recommendation system and algorithm centers around the “For You” page.

What is the For You feed?

The [For You feed](#) (a.k.a., For You page or FYP) is a curated stream of videos, a unique and tailored feed to the user’s interests.

Although TikTok users will see the same videos, your FYP is

totally unique and curated only for you.

Wenn Sie die App öffnen, landen Sie zuerst auf diesem Feed. Die gute Nachricht ist, dass Sie Ihren eigenen FYP kuratieren können, indem Sie mit der App interagieren.

Wenn Sie TikTok zum ersten Mal beitreten, wird Ihr FYP wahrscheinlich anhand der Interessen innerhalb Ihrer demografischen Gruppe basierend auf den von Ihnen angegebenen Informationen (dh Alter, Geschlecht und Interessen) kuratiert.

Es wird ein breiter Ansatz sein, bis Sie TikTok zeigen, wovon Sie mehr sehen möchten.

Werfen wir einen Blick auf die Aktionen, die den Algorithmus beeinflussen.

Wie funktioniert der TikTok-Algorithmus?

Die TikTok-Dokumentation nennt [drei Kernfaktoren, die den Algorithmus beeinflussen](#) :

- Benutzerinteraktionen.
- Videoinformationen.
- Geräte- und Kontoeinstellungen.

1. Faktor: Benutzerinteraktionen

Die Benutzerinteraktion umfasst Aktionen, die Benutzer für ein bestimmtes Video ausführen. Diese Wechselwirkungen können negativ oder positiv sein.

Eine positive Interaktion führt dazu, dass die App Ihnen mehr davon zeigt, und eine negative Interaktion bewirkt das Gegenteil. Mehr zu negativen Aktionen und ihren Auswirkungen auf „Was der Algorithmus nicht zeigt“.

Positive Interaktionen umfassen Engagements wie:

- Posten eines Kommentars.
- Einem Konto folgen.
- Liken eines Videos.
- Wiedergabezeit des Videos.

Wenn Sie als Ersteller wissen, wie der TikTok-Algorithmus funktioniert, wissen Sie, was Sie von Ihren Zuschauern verlangen müssen, um die richtige Art von Engagement zu steigern.

User interactions in practice

Benutzerinteraktionen sind sicherlich ein Faktor im Algorithmus, und TikToker wissen das.

Der TikTok-Algorithmus hat es richtig gemacht, als er vor ein paar Wochen das folgende virale Video auf meinem FYP platzierte. Ich beschäftige mich nicht sehr oft mit TikTok-Videos, aber das folgende Video hat mir gefallen und ich habe einen Kommentar darüber hinterlassen, wie gesund der Inhalt war. Ich wollte dem Schöpfer nicht folgen, aber ich genoss die kurze Interaktion.

https://www.tiktok.com/embed/v2/7140327363916696874?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxCcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSBZAKoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

Das Video sammelte 785.200 Links, 9060 Kommentare, 76.900 Speicherungen und 5.579 Shares. Am beeindruckendsten ist, dass das Video 6.000.000 Aufrufe hat.

Dank meines Likes und Kommentars landete ein zweites, ähnliches/Folgevideo des Erstellers auf meinem FYP.

Das zweite Video hatte nicht ganz die gleiche Magie, aber die Viralität von Video eins hat sicherlich dazu beigetragen, dass das zweite Video des Erstellers auf der Grundlage früherer Benutzerinteraktionen an den richtigen Stellen gelandet ist.

Zu Ehren von TikTok scheint ein gewisses Maß an Vertrauen in den TikTok-Algorithmus vorhanden zu sein.

Bei einigen der ansprechendsten Videos – oder Videos, die in einer Serie erstellt wurden – sehen Sie oft Kommentare wie „Vertraue dem Algorithmus, um mich zurückzubringen.“

TikToker könnten so etwas kommentieren, wenn sie dem Algorithmus zeigen möchten, dass sie an den Inhalten interessiert sind, aber (vielleicht) dem Ersteller nicht folgen möchten.

Im folgenden Video erinnert ein Ersteller die Zuschauer daran, dass sie auf „Folgen“ klicken können, anstatt dem Algorithmus zu vertrauen. Ein Folgen ist eine wünschenswerte Aktion für Ersteller, da es fast garantiert, dass der Benutzer mehr von ihren Inhalten sieht.

https://www.tiktok.com/embed/v2/7147085461222165803?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxCcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSBZAKoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

Wie man Benutzerinteraktionen fördert

Als Ersteller möchten Sie Benutzerinteraktionen fördern. Sie können dies erreichen, indem Sie großartige Inhalte erstellen, aber Sie können passive Zuschauer mit einem Aufruf zum Handeln

zu sinnvollen Aktionen anregen.

TikToker Tyla Brimblecombe von Styla Socials erklärt die Wichtigkeit eines Aufrufs zum Handeln in ihrem TikTok.

https://www.tiktok.com/embed/v2/7086951809331383553?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxCcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSBZAkoIG7r_uv4mhUd6fC2nH4cjIqlZo3w

Sie sagt:

- „Wir alle wissen, dass die wichtigste Metrik auf TikTok die Wiedergabezeit ist ... jemand scrollt durch den For You-Feed, er sieht Ihr Video, es gefällt ihm, er schafft es bis zum Ende. Wenn sie Ihren Aufruf zum Handeln sehen, „Schauen Sie sich meine anderen Videos an“, klicken sie auf Ihr Profil, sie sehen sich Ihre anderen Videos an, wodurch TikTok weiß, dass Ihre Inhalte wertvoll sind und Ihre Inhalte einer breiteren Öffentlichkeit zugänglich gemacht werden Publikum.“

2. Faktor: Videoinformationen

Laut TikTok-Dokumentation umfassen Videoinformationen:

- Geräusche
- Bildunterschriften
- Hashtags
- Videobeschreibungen
- Textüberlagerungen

Videoinformationen in der Praxis

TikTok ist bekannt für seine Assoziation mit trendigen Sounds. Mit einem Trendton kann der TikTok-Algorithmus eine Gruppe von Personen identifizieren, die sich wahrscheinlich mit einem Video beschäftigen, basierend auf früheren Benutzerinteraktionen. Wenn ein Benutzer jedes Mal, wenn ein bestimmter Ton verwendet wird, ein Video bis zum Ende ansieht, gefällt ihm der Inhalt wahrscheinlich und er möchte mehr davon sehen.

Es ist wichtig zu beachten, dass trendige Sounds nicht nur virale Tänze sind. Marken können Sounds verwenden, um schnell für ein breiteres Publikum sichtbar zu werden.

Kristyn Higginson, TikTok-Schöpferin von Skinician, sagt:

- „Die Erstellung ‚viraler‘ Inhalte schien in den Anfangstagen eine Herausforderung zu sein, da ein Großteil der Trendmusik nicht für Geschäftskonten verfügbar war. Der Ausweg bestand darin, sich für trendige Sounds zu entscheiden, wenn wir zuordenbare, relevante Inhalte für die Plattform erstellen wollten. Das bedeutete, dass wir uns als Marke ohne Urheberrechtsverletzung an Trends anlehnten. Dies könnte eine nützliche Taktik für andere kleine Marken in den frühen Phasen der Inhaltserstellung auf der Plattform sein.“

Mit der frühzeitigen Einführung eines angesagten Sounds erhielt Skincians viralstes Video 50.900 Aufrufe, 2.142 Likes und 40 engagierte Kommentare. Das Video stellte die Marke ihrem TikTok-Publikum vor.

https://www.tiktok.com/embed/v2/7090277499564772614?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-

m5Hu_2r8qYWG05wFxKcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSB
ZAkoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

TikTok-Ersteller stellen fest, dass Videoinformationen erheblich zum Videoerfolg beitragen.

Kate Smoothy von Web Hive Digital ist eine TikTok-Erstellerin.

Smoothy teilt mit, dass Bildunterschriften, Hashtags und Videobeschreibungen ein zentraler Bestandteil ihrer Strategie sind. Sie sagt:

- „Ich kann gar nicht genug ausdrücken, wie sehr Sie Ihre Videos mit Untertiteln versehen müssen. Dies ist wichtig für diejenigen, die nichts hören, aber für diejenigen, die es vorziehen, ohne Ton durch TikTok zu scrollen. Sie werden Ihre Zielgruppengröße und -reichweite ohne Untertitel erheblich reduzieren. Es gab noch nie einen besseren Zeitpunkt, Schlüsselwörter in Ihre TikTok-Beschreibungen aufzunehmen. Wir sehen eine große Veränderung bei der Plattform, wo Videos, die für ihre Suchfunktion optimiert sind, eine bessere Leistung zu erbringen scheinen.“

Smoothy schreibt den Erfolg ihres „Super-Low-Effort-Marketing-Hack-Videos“ der Aufnahme von Video-Informationstaktiken zu.

https://www.tiktok.com/embed/v2/7134650653049130245?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxKcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSBZAkoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

Das Video hat 92.400 Aufrufe und 5.913 Likes, 91 Kommentare

und 740 Shares.

Die in der Videobeschreibung, innerhalb des Text-Overlays und innerhalb der Videobeschreibung verwendeten Keywords haben TikTok mitgeteilt, worum es in dem Video geht. Dies führte zu einem „Ranking“ des Videos für den Suchbegriff „Marketing-Hack“ und „Marketing-Tipps“ innerhalb der Suchfunktion von TikTok.

So verwenden Sie Videoinformationen

Der Top-Tipp von TikTok-Experten ist, eine Schlüsselwortrecherche mit der TikTok-App durchzuführen und dann Schlüsselwörter in Ihrem Video, Ihren Bildunterschriften und Hashtags zu verwenden.

https://www.tiktok.com/embed/v2/7139996958801825070?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxKcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSBZAkoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

3. Faktor: Geräte- und Kontoeinstellungen

TikTok ist klar, dass Geräte- und Kontoeinstellungen ein geringeres Gewicht im Algorithmus erhalten.

Zu den offensichtlicheren Faktoren gehören Sprachpräferenzen und Ländereinstellungen. Es ist sinnvoll, dass Benutzer Videos in der Sprache sehen möchten, die sie sprechen.

TikTok bezieht auch den Gerätetyp in seinen Algorithmus ein. Sie ziehen es möglicherweise vor, Ihnen Videos zu zeigen, die Ihr Gerät nahtlos abspielen kann, da dies zu einer positiven Benutzererfahrung für Personen führen würde, die ältere Telefone oder kleinere Bildschirme verwenden.

Keywords und der TikTok-Algorithmus

In den letzten Monaten wurde viel von TikTokers über eine Änderung des TikTok-Algorithmus und darüber gesprochen, wie TikTok jetzt Schlüsselwörter priorisiert.

Wenn Sie nach „ [TikTok Algorithm Change](#) “ suchen, werden Sie Leute finden, die viel das Gleiche sagen – TikTok SEO wird immer häufiger.

https://www.tiktok.com/embed/v2/7137660935174458666?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxFKcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSBZAKoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

Interessanterweise scheint TikTok nirgendwo in seiner Dokumentation „Schlüsselwörter“ in Bezug auf seinen Algorithmus zu erwähnen.

Aber das bedeutet nicht, dass Schlüsselwörter nicht wichtig sind.

Es ist kein Zufall, dass [TikTok die Videobeschreibungen auf 2.200 Zeichen erhöht hat](#) . In ihrer Ankündigung über die App sagte TikTok:

„Mehr Charaktere geben den Erstellern die Möglichkeit, Engagement zu generieren, während sie besser durchsuchbar und von TikTok den Zuschauern besser empfohlen werden.“

TikTok erwähnt auch [Bildunterschriften und Hashtags](#) , die beide zu FYP beitragen.

Bildunterschriften und Hashtags sind **Wörter** , die dabei helfen, **Videos in einen Kontext** zu setzen . Wenn Untertitel

ein Faktor sind, deutet dies darauf hin, dass Wörter – oder Schlüsselwörter – zum TikTok-Algorithmus beitragen.

Wir wissen auch, dass TikTok-Benutzer Inhalte basierend auf Schlüsselwörtern einschränken können (dazu später mehr).

TikTok-SEO und Google

Um TikTok SEO weiter zu erkunden und wie sich dies auf der weltweit beliebtesten Suchplattform Google abspielt. Wir haben einige Tests durchgeführt, um die Indexierung von TikTok-Videos durch Google zu untersuchen.

Wenn Sie bei Google nach „ [TikTok-Rezepten](#) “ suchen, erwarten Sie TikTok-Videos. Aber was wir fanden, war das Gegenteil.

Google priorisierte Websites, die die viralen Rezepttrends von TikTok teilen. TikTok rangiert mit einer Tag-Seite auf der achten Suchposition. In Anbetracht der Relevanz zum Suchbegriff ist dies kein allzu starker Rang.

In den Video-Ergebnissen belegte TikTok den sechsten Platz unter fünf YouTube-Videos. Es macht Sinn, dass Google seinen eigenen Kanal YouTube über TikTok priorisiert.

Es ist erwähnenswert, dass Videos, die im TikTok-Rezept-Tag enthalten sind, alle das Hashtag #recipes in der Bildunterschrift enthielten. Bildunterschriften helfen dem TikTok-Algorithmus, Videos nach Themen zu sortieren, und könnten sogar die Sichtbarkeit über die Google-Suche verbessern.



Wie im Screenshot gezeigt, enthält die Bildunterschrift das Hashtag #recipes, was darauf hindeutet, dass Schlüsselwörter für den TikTok-Algorithmus relevant sind.

Also, was bedeutet das?

- Ersteller sollten relevante Schlüsselwörter in ihren Bildunterschriften und Hashtags verwenden.
- Laut TikTok-Dokumentation tragen Bildunterschriften und Hashtags zum TikTok-Algorithmus bei.
- Erwägen Sie für eine bessere Sichtbarkeit in Google, Ihr Video auf anderen Plattformen wie YouTube Shorts zu teilen.

Holen Sie sich den täglichen Newsletter, auf den sich Suchmaschinenvermarkter verlassen.

[Siehe Bedingungen.](#)

TikTok-Algorithmus-Mythen entlarvt

Lassen Sie uns mit einigen weit verbreiteten Missverständnissen über TikTok aufräumen.

Mythos: Du musst täglich posten

Das Posten von 1-3 TikTok-Videos täglich ist eine häufige Empfehlung für neue TikToker. Die Richtlinien von TikTok stellen jedoch klar, dass [nicht erforderlich dies für das Wachstum](#) ist .

Allerdings kann es sinnvoll sein, [verschiedene Videos zu testen](#) . Die Dokumentation von TikTok weist darauf hin, dass das Experimentieren mit Videos, die Sie teilen, und das Posten hochwertiger Inhalte nützlicher ist, um ein Publikum zu beschäftigen.

Mythos: Creators im TikTok Creator Fund erhalten mehr Auffindbarkeit

Die gute Nachricht ist: Du musst nicht im Creator Fund sein, um deine Chancen auf Auffindbarkeit zu erhöhen.

Diejenigen im TikTok Creator Fund werden mit größerer Wahrscheinlichkeit mehr Engagement erzielen, aber das liegt daran, dass diese TikToker wissen, wie man großartige Inhalte erstellt, die die Community ansprechen.

Mythos: TikTok ist nur für Kurzformvideos

Obwohl TikTok als App mit sieben Sekunden langen Videos gestartet ist, hat sich die maximale Videolänge seitdem erhöht.

Innerhalb der App können Sie Videos mit einer Länge von 15 Sekunden, 60 Sekunden oder drei Minuten erstellen.

Alternativ können Sie erweiterte Videos erstellen und in die TikTok-App hochladen.

TikTok erwähnt nicht, dass längere Videos im Algorithmus beliebt sind, aber sie scheinen Gründe zu teilen, [warum längere Videos](#) mit Erstellern interagieren.

https://www.tiktok.com/embed/v2/7068672269693848838?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxKcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSBZAKoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

Was der Algorithmus nicht zeigt

Um TikTok Ehre zu machen, ergreift es Maßnahmen, [um die Empfehlungen](#) im FYP zu schützen und zu diversifizieren.

Mit einem Empfehlungssystem wäre es für TikTok leicht, sich zu wiederholen und nur ähnliche Videos zu zeigen. Ebenso könnte es für TikTok einfach sein, seinen Zuschauern unerwünschte Inhalte zu zeigen.

TikTok ergreift Maßnahmen, um TikTok zu einem besseren Ort für alle zu machen. Folgendes zeigt TikTok im FYP nicht an.

Doppelte Inhalte

TikTok sagt, dass zu viel von allem ermüdend werden kann.

Aus diesem Grund zeigt [TikTok im FYP keine doppelten Inhalte an](#). Außerdem sorgt die Diversifizierung des FYP dafür, dass die Benutzer einer Reihe von Ideen und Perspektiven ausgesetzt sind.

Potenziell problematische Cluster

Der TikTok-Algorithmus erfüllt [die Mission der Plattform](#), Kreativität zu inspirieren und Freude zu bereiten.

Um dies zu erreichen, vermeidet TikTok, ähnliche Inhalte zu potenziell problematischen Themen wie extreme Diäten oder Fitness, Traurigkeit oder Trennungen zu empfehlen.

TikTok möchte seine Nutzer davor schützen, Inhalte anzusehen, die bei gemeinsamer Betrachtung Stress verursachen können, aber als einzelnes Video in Ordnung sein könnten.

Diese Art der Einschränkung wird durch die Funktionen „Kein Interesse“ und „Keyword-Filterung“ positiv verstärkt.

Als „Kein Interesse“ gekennzeichneteter Inhalt

Mit TikTok können Benutzer steuern, was sie auf der Plattform sehen möchten oder nicht sehen möchten.

Zuschauer können bei jedem TikTok-Video den Bildschirm gedrückt halten und eine der folgenden Optionen auswählen:

- **Video speichern:** Ein positiver Video-Engagement-Indikator für den Algorithmus.
- **Löschmodus** : Videos ohne Benutzernamen und Untertitel usw. anzeigen.
- **Melden** : Diese Aktion weist darauf hin, dass ein Video gegen die Community-Richtlinien von TikTok verstößt und behandelt werden muss.
- **Kein Interesse** : Dies teilt dem TikTok-Algorithmus mit, dass Ihnen als Zuschauer ein bestimmter Inhalt nicht gefällt. Im Gegenzug sieht man weniger davon.

↓ Save video

☑ Clear mode

🚩 Report

💔 Not interested

Wenn ein Zuschauer auf „Kein Interesse“ klickt, können Sie

damit rechnen, dass dies den Algorithmus und seinen FYP beeinflusst.

Inhalt einschließlich gefilterter Schlüsselwörter

Benutzer können TikTok mitteilen, was sie nicht sehen möchten, indem sie zum Menü oben rechts gehen und zu *Einstellungen und Datenschutz > Inhaltseinstellungen > Videoschlüsselwörter* filtern navigieren .

Hier können Sie Schlüsselwörter hinzufügen und Videos von Benutzern, denen Sie folgen, und/oder dem FYP filtern.

14:53



Zoe Ashbridge Freel... ▾



@zoefreelanceseo



31

Following

860

+1

Followers

3090

Likes

Edit profile

Expert-Vetted Freelancer

Daily SEO & freelance tips to land clients

Follow 🚀



My orders



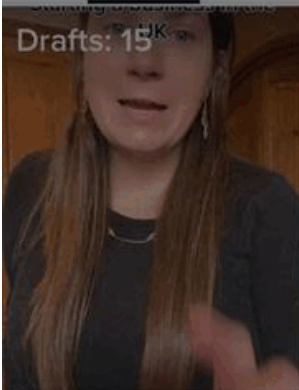
Q&A



Supporting: Cardiac Risk in the



Drafts: 15



Pinned



Creator tools



My QR code



Settings and privacy

Inhalte, die von Personen unter 18 Jahren erstellt wurden

In seinem Bestreben, die App zu einem sicheren Ort für alle zu machen, hat TikTok [Maßnahmen](#) zum Schutz seiner jüngeren Nutzer eingeführt.

Kinder oder Jugendliche unter 18 Jahren können keine Inhalte im Explore-Feed präsentieren – was bedeutet, dass ihre Inhalte und Profile für andere TikTok-Benutzer nicht so leicht zu finden sind.

7 Tipps, um mit dem TikTok-Algorithmus zu arbeiten und viral zu werden

Einige Top-TikTok-Ersteller bieten einige Best Practices für die Erstellung von Videoinhalten, die die Gunst des Algorithmus gewinnen könnten.

Tipp 1: Wechseln Sie zu einem TikTok Pro-Konto

[TikTok Pro-Konten](#) bieten Erstellern detailliertere Analysen.

Wenn Sie eines erstellen, können Sie Erkenntnisse zu Folgendem sehen:

- Wöchentliche und monatliche Ansichten.
- Follower-Wachstum.
- Trendvideos.

Wenn Sie ein Liebhaber eingehender Analysen sind, ist dies möglicherweise das Richtige für Sie.

TikTok davon abräät [Es ist jedoch wichtig zu beachten, dass die](#)

[Anleitung von](#) , Inhalte hauptsächlich rund um Analysen zu erstellen. Denken Sie stattdessen beim Erstellen von Inhalten an das Gesamtbild.

Tipp 2: Finden Sie Ihre Nische

Das Finden einer Nische auf TikTok gibt Ihnen ein besseres Verständnis für das Gesamtbild.

Wenn Sie wissen, in welche Nische Sie fallen, können Sie Themen mithilfe der Suchleiste erkunden und genau sehen, welche Videos gut abschneiden, damit Sie so etwas wie diese nachbauen können.

https://www.tiktok.com/embed/v2/7043578829515803950?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxFKcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSBZAkoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

Tipp 3: Verwenden Sie die richtigen Hashtags, fügen Sie Bildunterschriften hinzu und schreiben Sie Videobeschreibungen

Keywords haben eindeutig Vorteile, wie oben erwähnt.

Denken Sie daran, sich für Hashtags zu entscheiden, die für Ihr Video oder Ihre Nische sehr relevant sind, um eine Chance zu haben, in den Top-Listen für Keywords zu erscheinen.

https://www.tiktok.com/embed/v2/7136234072271047942?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-

NDc_-

m5Hu_2r8qYWG05wFxFcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSB
ZAkoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

Tipp 4: Verwenden Sie trendige Sounds und Musik

Egal wie sehr sich der Algorithmus ändert, Musik und trendige Sounds werden wahrscheinlich einen Platz auf der TikTok-Plattform haben.

Denken Sie daran, dass es bei trendigen Sounds und Musik nicht immer um Voiceovers und virale Tanzbewegungen geht.

Sie können sie verwenden, um Ihre Marke vorzustellen, sie in Ihrer Nische identifizierbar zu machen und sich abzuheben.

Tipp 5: Erstellen Sie qualitativ hochwertige Videos

[Mit den Bearbeitungstools von TikTok](#) können Benutzer auffällige, qualitativ hochwertige Videos innerhalb der App erstellen und bearbeiten.

TikTok-Ersteller betonen, wie wichtig es ist, frühzeitig Engagement zu schaffen. Versuchen Sie also, Ihre Videos so zu bearbeiten, dass die Aufmerksamkeit des Benutzers in den ersten drei Sekunden auf sich gezogen wird.

https://www.tiktok.com/embed/v2/7003425150288792838?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxFcphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSBZAkoiG7r_uv4mhUd6fC2nH4cjIqlZo3w

Tipp 6: Tauschen Sie sich mit anderen TikTok-Benutzern aus

Dein Kommentarbereich ist eine Fundgrube für Videoideen! Und Sie können auf Kommentare mit einem Video antworten.

Diese Funktion stellt Ihre Antwort im Videoformat in den Kommentarbereich.

Sie können darauf wetten, dass andere Kommentatoren die gleichen Fragen haben, sodass sie sich wahrscheinlich durchklicken und bei Ihrem nächsten Video mitmachen.

Tipp 7: Veröffentlichen Sie eine Serie

Etwas im Gegensatz zum Langform-Videoformat können kürzere Videos in Serie immer noch effektiv sein.

Henry Purchase von A Couple Things To Do hat eine TikTok-Fangemeinde auf 193.000 Follower anwachsen lassen. Kaufen Sie Aktien, Ziel und Aktivitäten für Paare. Er sagt:

- „Wir posten mehrere Videos derselben Aktivität und decken jedes Mal mehr Informationen auf – das hält die Leute für mehr zurück. Zum Beispiel haben wir zunächst nicht den Standort unseres beliebtesten Videos gepostet. Nachdem es an Popularität gewonnen hatte, nutzten wir Kommentare, um weitere Inhalte zu erstellen. Daher die Wiederverwendung von Inhalten und die Interaktion mit unserem Publikum.“

https://www.tiktok.com/embed/v2/7162637530427346182?lang=de&referrer=https%3A%2F%2Fsearchengineland.com%2Fhow-tiktok-algorithm-works-390229%3Fmkt_tok%3DNzI3LVpRRS0wNDQAAAGI2EEqlvCBB3e-NDc_-m5Hu_2r8qYWG05wFxFkphvCbR9onJuW2wHHRf525KCoCttrxMMRt4SP_FskSB

Warum sich Marken für TikTok interessieren sollten

Lieben Sie es oder hassen Sie es, TikTok hat sich bewährt, da sein großzügiger Algorithmus sich gut für schnelles Wachstum eignet.

Mit ein wenig Engagement können Marken in kürzester Zeit Tausende von Menschen erreichen.

Obwohl qualitativ hochwertige Videos bevorzugt werden, können einfache Videos vor der Kamera genauso viel Aufmerksamkeit erregen.