

Modulo di istruzione di InCyT

Sicurezza delle informazioni per i dipendenti

Unità 2 Settimana 3

Nome dell'unità: Malware e attacchi tipici

<i>Attività di apprendimento</i>	☒ Webinar ☐ Exercises ☐ Workshop ☐ Presentation ☐ Other
<i>Responsabile dell'apprendimento</i>	Claudio Fornaro
<i>Obiettivi dell'attività di apprendimento</i>	1. Introduzione ai vari tipi di malware 2. Conoscenza dei tipi più comuni di attacchi
<i>Competenze da acquisire</i>	• Conoscenze di base su cos'è il malware e gli attacchi più comuni ai sistemi informatici • Soft Skill 2 - SS2
<i>Durata dell'attività di apprendimento</i>	1 settimana
<i>Requisiti di apprendimento</i>	Cosa è richiesto • Non è richiesta alcuna conoscenza specifica precedente

Brevi informazioni sul modulo

Chiamandolo di malware, il nome implica già qualcosa di dannoso. Questo termine si riferisce generalmente a software progettati intenzionalmente per causare danni di varia natura, in particolare (ma non solo): il blocco di servizi e comunicazioni, l'intercettazione di trasmissioni, il furto di informazioni private (personali, aziendali o governative) e di identità personali, l'accesso non autorizzato a informazioni o sistemi, la negazione agli utenti dell'accesso alle informazioni.

Dopo una parte introduttiva verranno delineati i concetti fondamentali legati al malware e descritti i tipi principali, i rischi che comportano e come possano essere mitigati.

La seconda parte del modulo intende offrire una panoramica degli attacchi più comuni che un sistema può subire, elenco non completo poiché ogni giorno ne vengono scoperti nuovi tipi. Questi attacchi sono tentativi di accedere a un sistema informatico o di interromperne i servizi. La descrizione parte dagli attacchi più comuni e semplici (Phishing) per arrivare a quelli un po' più complessi: attacchi alle password, buffer overflow, SQL injection, cross-site scripting (XSS), cross-site request forgery (CSRF), attacchi relativi ai file, denial of service (DoS) e distributed denial of service (DDoS), man-in-the-middle (MITM), ARP spoofing o ARP poisoning, DNS spoofing, attacchi guidati (drive-by), riutilizzo delle credenziali, dirottamento di sessioni Web, dirottamento di sessioni TCP, spoofing IP, exploit zero-day e conseguenze della perdita di codice.

Contenuto del materiale didattico

In informatica, "hacker" e "cracker" sono due termini che vengono utilizzati per indicare persone con una grande conoscenza dei sistemi informatici. La principale differenza tra i due sta nell'atteggiamento con cui agiscono. Un **hacker** (il termine nasce nella seconda metà del 1900 al MIT) agisce esplorando i sistemi per capire meglio come funzionano e cercando di scoprire nuove falle per risolverle e quindi aumentare la sicurezza di un sistema. Alcune azioni possono essere solo dimostrative o per ragioni etiche, ma mai a scopo di lucro. Un **cracker**, d'altra parte, usa le sue abilità cercando illegalmente di entrare nei sistemi per trarne profitto rubando dati o in modo distruttivo.

Gli hacker cercano costantemente di acquisire nuove conoscenze sui sistemi informatici per comprenderli meglio o migliorarne le funzioni, le prestazioni e la sicurezza. Alcuni di loro si impegnano anche in azioni dimostrative ed etiche, ma mai a scopo di lucro. Le azioni degli hacker vengono eseguite quasi sempre senza alcuna autorizzazione, quindi la legge le punisce (con grandi differenze tra i Paesi), per questo motivo gli hacker di solito usano soprannomi. A volte, le società di sicurezza o IT invece di far perseguire gli hacker, spesso li assumono per sfruttare le loro elevate capacità e creatività (eventualmente dopo che questi hanno pagato il loro debito con la Giustizia). Vengono quindi impiegati per fare il test di applicazioni, sistemi di sicurezza e programmi sia in produzione che già sul mercato al fine di bloccare gli attacchi e prevenire problemi che potrebbero

sorgere in futuro da bug dei sistemi.

I *veri* hacker hanno una reputazione da esperti e sono molto rispettati nelle loro comunità. Vengono organizzate gare con l'obiettivo di "bucare" un sistema o un programma nel più breve tempo possibile. Queste competizioni sono chiamate hackathon e sono spesso sponsorizzate dai giganti dell'industria informatica.

Il termine "cracker" è attribuito a Richard Stallman (una figura ben nota del movimento del software libero, ideatore del progetto GNU e membro della comunità hacker sin dagli anni '70). Questi ha cercato di evitare l'abuso del termine "hacker", differenziando chi agisce senza alcuna volontà di nuocere e chi, al contrario, agisce per trarne profitto o solo per arrecare danno. Un cracker è un cyber-criminale che tenta di aggredire i sistemi informatici in modo non autorizzato e illegale, con malevolenza, cercando di trarne vantaggio fraudolento sfruttando vulnerabilità non ancora conosciute e risolte. I cracker sono spesso coinvolti in spionaggio industriale, furto di dati e blocco di reti e siti Web, atti di vandalismo, ecc.

I termini hacker e cracker hanno limiti sottili che li dividono. Ad esempio un'intrusione può essere considerata un atto illegale da alcuni e non da altri e la Legge rispecchia queste opinioni; ma questo documento non riguarda aspetti etici o legislativi. Bisogna invece considerare che oggi, purtroppo, i due termini sono considerati sinonimi.

In questo documento, useremo contro voglia il più comune e abusato termine hacker per entrambe le categorie, ben sapendo che l'accorto lettore sarà perfettamente in grado di distinguere tra un hacker e un cracker.

Gli "hacker" possono essere classificati grossolanamente in diverse categorie, le più utilizzate sono le seguenti.

I **black hat** hacker utilizzano le loro conoscenze per scopi malevoli o criminali, facendo attenzione a non essere rintracciati.

I **white hat** hacker testano i sistemi per identificare le possibili vulnerabilità ed eliminarle per evitare accessi indesiderati. Spesso si tratta di esperti di sicurezza pagati a questo scopo.

I **gray hat** hacker sono simili ai white hat, ma spesso conducono ricerche sulle vulnerabilità del sistema senza autorizzazione e, una volta scoperte, le utilizzano per costringere i fornitori a risolvere il problema, di solito gratuitamente. Alcuni fornitori li ricompensano con denaro, altri li perseguono.

I **suicide** hacker (suicidi) sono di solito meno competenti dei black/white hat, ma agiscono per motivi ideologici e non hanno interesse a nascondere la propria identità.

Gli **sponsored** hacker sono commissionati dallo Stato, sono pagati da un governo per sferrare attacchi informatici contro un altro Paese, hanno accesso a grandi risorse fornite dallo Stato e

sono protetti dal perseguimento dei reati commessi. Agiscono per sottrarre informazioni riservate, piani o progetti, o per creare disordini generalizzati e danneggiare la popolazione, l'esercito, la gestione del governo, gli ospedali, gli impianti per la produzione di energia, l'acqua e l'alimentazione, ecc. Queste azioni sono considerate vere e proprie azioni di guerra.

Gli **script kiddies** sono hackers che non hanno o non hanno ancora acquisito grandi capacità, di solito utilizzano script e strumenti creati da altri hacker. Tutti gli hacker prima di acquisire grandi competenze sono script kiddies.

I **terroristi** informatici sono individui o gruppi con intenzioni ostili che generalmente mirano a creare il caos in un Paese; sono simili agli hacker commissionati da uno Stato, ma non hanno (ancora) un Governo che li sostenga, o il sostegno non è evidente.

Malware

Malware è l'acronimo di "software maligno", un termine generico che identifica qualsiasi programma dannoso per un sistema. Il malware è il software che gli hacker creano per perseguire i loro obiettivi: ricavare soldi, rubare dati, interrompere le attività, ecc.

Uno dei malware più comuni è il cosiddetto **trojan horse** (cavallo di Troia, anche semplicemente *Trojan*), un programma che sembra essere un software utile e spesso si comporta come tale, ma che porta con sé una funzionalità nascosta che viene eseguita insieme a quella normalmente attivata all'avvio dell'applicazione. Spesso si tratta di programmi shareware o di programmi sprotetti illegalmente ("cracked") che vengono distribuiti in forma modificata al fine di veicolare il cavallo di Troia. In altri casi si tratta solo di malware con un nome falso che lo fa sembrare un software utile. I trojan vengono spesso diffusi tramite e-mail o scaricando programmi dal Web e di solito non si diffondono da soli ad altri computer. Esempi di malware sono i *keylogger* (programmi che rubano le password e le inviano all'hacker), le *backdoor* (programmi che funzionano come terminal server e permettono all'hacker remoto di accedere al computer in qualche modo aggirando le normali procedure di autenticazione), i *cryptominer* (programmi che estraggono criptovalute), i *bot* (abbreviazione di robot, usati ad esempio per un attacco distributed denial of service).

I **rootkit** sono software installati da malware che cercano di nascondere la presenza di un malware stesso al rilevamento; ottengono questo risultato modificando il sistema operativo.

Il malware che infetta altri software può essere un virus o un worm. Questi ultimi prendono il nome dal modo in cui si diffondono piuttosto che da uno specifico tipo di comportamento. Un **virus** informatico è un software simile al cavallo di Troia, ma può produrre copie di se stesso infettando (introducendosi) altri software, compresi i file del sistema operativo. Il virus si diffonde ad altri sistemi quando il software infetto viene copiato ed eseguito su questi sistemi.

Simile a un virus è il **worm**, un software autonomo che opera su un computer e infetta altri

computer attraverso la rete, ma senza infettare i file.

Un **ransomware** è un tipo di malware che limita l'accesso al dispositivo che infetta, richiedendo il pagamento di un riscatto per rimuovere la restrizione. I principali tipi di ransomware sono 1) *Screen-locking ransomware* che blocca gli schermi con un'accusa intimidatoria di comportamenti illeciti e chiedendo alle vittime di pagare un compenso; 2) *Encryption-based ransomware* in cui il software dannoso cifra alcuni tipi di file (immagini, documenti) o tutti i file su una macchina infetta. Al termine della crittografia completa, un pop-up di solito informa l'utente che (tutti) i file sono stati crittografati e che è necessario pagare una tariffa per recuperarli, di solito in criptovaluta. Esempi di encryption-based ransomware sono CryptoLocker e WannaCry. Una variante nota come **Wiper** sembra un ransomware, ma distrugge comunque i dati, anche se spesso è richiesto un pagamento.

I **Greyware** sono una categoria diversa di malware: questi programmi possono peggiorare le prestazioni dei computer e causare rischi per la sicurezza, comportarsi in modo fastidioso o indesiderato, ma di solito sono meno pericolosi (che è un'opinione relativa) rispetto alle precedenti categorie di malware. Esempi di greyware sono: *Spyware* (es. keylogger, programmi che registrano tutti i tasti premuti sulla tastiera), *adware* (finestre pubblicitarie sullo schermo o nel browser, dialer fraudolenti, ecc.), *Programmi Potenzialmente Indesiderati* (PUP, applicazioni che l'utente ha installato solo a causa di un errore in una finestra di accordo durante l'installazione di un altro programma). Possono essere programmi autonomi o plug-in per altri programmi, ad esempio i browser. A volte si limitano a modificare le impostazioni di altri programmi, ad esempio per portare l'utente a un sito Web pubblicitario o mostrargli una pagina di annunci ogni volta che viene avviato il browser o occasionalmente.

Poiché i motori antivirus potenziano il loro potere di rilevamento del malware, viene utilizzata una combinazione di varie tecniche per cercare di evitare il rilevamento e l'analisi. I malware più recenti sono in grado di adottare sofisticate contromisure:

- analizzano e rilevano tramite fingerprinting l'ambiente per adottare specifiche modalità di occultamento;
- cercano di offuscare se stessi per confondere il rilevamento antivirus in base alla firma del malware;
- realizzano un'evasione temporale per ingannare il rilevamento di comportamenti anomali, il malware è normalmente quiescente e si attiva in periodi molto specifici (ad esempio durante il processo di avvio) o a seguito di determinate azioni intraprese dall'utente.

Una tecnica di evasione molto difficile da scoprire è il *Fileless malware* o *Advanced Volatile Threats*. Questo tipo di malware viene eseguito in memoria e non richiede un file per funzionare. Utilizza gli strumenti di sistema esistenti per svolgere la propria attività. L'assenza di qualsiasi file

nel file system li rende piuttosto difficili da rilevare e analizzare, ad eccezione, in parte, di quegli strumenti di monitoraggio che controllano l'utilizzo delle risorse e il comportamento del sistema. L'aumento di questo tipo di attacchi è molto alto.

Rischi

Il software vulnerabile può essere sfruttato per ottenere privilegi o aggirare le difese fino a quando il problema non è stato risolto; questo è il motivo principale per cui il software deve essere aggiornato, in particolare il sistema operativo. Firewall e sistemi di rilevamento delle intrusioni monitorano i modelli di traffico sulla LAN.

La regola del privilegio minimo dice che utenti e programmi non dovrebbero avere più privilegi di quelli richiesti, altrimenti il malware può trarne vantaggio.

Mitigazione

Il software *antivirus* blocca ed eventualmente rimuove alcuni tipi di malware, possono fornire protezione in tempo reale contro l'installazione di software malware su un computer e una scansione completa di file e configurazioni (ad es. il Registro di Windows) alla ricerca di potenziali minacce.

I *firewall personali* analizzano qualsiasi connessione di rete e possono individuare attività insolite.

Il *sandboxing* consiste nell'eseguire un'applicazione isolata dal resto del sistema, alcune applicazioni di sandboxing possono eseguire un'applicazione specifica mentre in sistemi più complessi una macchina virtuale potrebbe essere più appropriata.

Le scansioni delle vulnerabilità del sito Web vengono eseguite da alcuni software antivirus su un elenco di siti pericolosi noti e impediscono la connessione ad essi.

La network segregation consiste nel dividere una rete in parti e abilitare solo quel traffico tra di esse noto per essere legittimo, questo impedisce la diffusione di software dannoso nella rete. Le moderne tecniche di Software Defined Networking sono in grado di realizzare questa segregazione.

L'isolamento "*Air gap*" o "*parallel network*" porta la segregazione della rete al limite isolando completamente una parte di una rete e fornendo controlli avanzati sull'ingresso e l'uscita di software e dati dal mondo esterno. Anche questa soluzione non è quella definitiva, in quanto software e dati devono entrare nel sistema isolato, ad esempio per il patching. Un noto esempio di violazione di una rete isolata è il famoso malware **Stuxnet** che è stato introdotto nell'ambiente di destinazione tramite unità USB dimenticate "casualmente" nelle vicinanze di un impianto,

prelevate dal personale e inserite nei computer dell'impianto per vederne il contenuto. Stuxnet ha tre moduli: un worm che controlla l'attacco; un file di collegamento che esegue automaticamente il worm quando viene inserito in un computer Windows; e un componente rootkit responsabile di nascondere tutti i file e i processi dannosi, per impedirne il rilevamento. Altri modi per attraversare gli air gaps analizzano le emissioni elettromagnetiche, termiche e acustiche.

Tipi di attacchi più comuni

Questa parte descrive alcuni degli attacchi più comuni a un sistema informatico. Non è e non può essere completo in quanto la varietà di attacchi è estremamente vasta, ognuno ha molte varianti e molti attacchi e varianti vengono creati e scoperti (e risolti) ogni giorno. Ciò significa che lo scopo di questa parte non è quello di fornire una classificazione completa, ma di dare un'idea dell'alto livello di conoscenza richiesto sia per perpetrare attacchi sia per contrattaccare.

Phishing

Gli attacchi di phishing sono molto comuni e facili da eseguire. Consistono nell'invio di e-mail che sembrano provenire da un mittente affidabile, come un collega o una banca, con l'obiettivo di indurre la vittima a compiere un certo tipo di operazione.

L'attaccante cerca di persuadere la vittima ad aprire un allegato o fornisce un link che porta ad una pagina molto simile a quella originale, ad esempio una home page di banca, ma creata dall'attaccante, per poi chiedere alla vittima di inserire le proprie credenziali. Per ottimizzare questo tipo di attacco ci sono tecniche che fanno parte della cosiddetta *social engineering*: in questo caso la vulnerabilità del sistema informatico è da individuare nell'utente stesso che viene ingannato e involontariamente fornisce le sue credenziali.

Per difendersi da questi attacchi, è necessaria un'attenta ispezione del link (ad es. è difficile distinguere tra la L minuscola 'l' e la I maiuscola o cifra uno '1', o notare la differenza tra .it e .It in un URL) . Inoltre, si dovrebbe essere consapevoli del fatto che le aziende serie non inseriscono mai collegamenti nelle loro e-mail atte a portare l'utente a una pagina di accesso.

Attacchi alle password

Per venire a conoscenza di una password si può provare con il social engineering, come visto nel paragrafo precedente, oppure si considera che queste vengono solitamente archiviate come hash che sono molto difficili da invertire. I casi sono essenzialmente due: la password crittografata è disponibile o è sconosciuta e questo tipo di attacco può essere effettuato in diversi modi:

- **Attacco a forza bruta:** in questo caso il programma proverà tutte le possibili combinazioni di personaggi di un set;
- **Attacco con dizionario:** il programma verifica una sequenza di password comunemente utilizzate (dizionario);
- **Attacco con file rainbow:** una tabella rainbow è una tabella precompilata che elenca i risultati di un attacco con dizionario. Questo accelera la ricerca. L'uso di una chiave calcolata mediante un seme iniziatore rende questo attacco impraticabile.

I programmi in grado di eseguire attacchi alle password, sia a forza bruta che a dizionario, dovrebbero essere utilizzati regolarmente dagli amministratori di sistema per controllare i file delle password dei loro utenti.

Per proteggersi da questi attacchi, un sistema dovrebbe verificare la qualità della password prima di permettere all'utente di impostarla, imponendo l'uso di un numero minimo di caratteri (ad esempio 10) provenienti da diversi set (lettere maiuscole e minuscole, cifre, caratteri di punteggiatura). Se la password immessa è non corretta per alcuni tentativi successivi, è appropriato adottare una politica di blocco dell'account o di ritardo al tentativo successivo, impedendo così all'attaccante di provare velocemente molte combinazioni.

Buffer Overflow

Il buffer overflow si verifica quando si inviano a un buffer più byte di quanti ne possa contenere, i byte traboccano dal buffer destinato a contenerli e modificano altri byte o, a seconda del sistema operativo e del tipo di file, addirittura iniettano codice macchina eseguibile. Mentre i sistemi operativi moderni possono impedire l'esecuzione di un file eseguibile con dati e codice misti, ciò non è vero per i sistemi operativi più semplici spesso presenti nei sistemi embedded, dove in molti casi non esiste alcun sistema operativo.

SQL Injection

Un attacco molto comune è la SQL Injection. Il linguaggio SQL è utilizzato nei database relazionali; pertanto questi tipi di attacchi sono finalizzati al furto di informazioni. In sostanza, data la natura interpretativa della query SQL, un codice SQL dannoso viene inviato a una pagina web che si sa che utilizza una query SQL per raccogliere i dati da presentare all'utente. Il codice SQL dannoso modifica (ad es. aggiungendo solo altre parti sintattiche) la query SQL predefinita ottenendo una query modificata che espone dati diversi da quelli per cui era stata scritta la query originale.

Sebbene la prevenzione di questo tipo di attacchi non sia complessa (ad esempio impedire l'accesso ad altre tabelle del database oltre a quelle necessarie per eseguire la query originale), la

SQL Injection è ancora uno degli attacchi più comuni.

Cross-site Scripting (XSS)

In un cross-site scripting (XSS) gli *script eseguibili* malevoli vengono iniettati nel codice di un'applicazione Web (plugin) attendibile o di una pagina del sito Web, con il risultato che viene scaricato da altre persone (visualizzatori). Quando il server restituisce la pagina Web compromessa, il browser Web dell'utente considera che sia stata inviata da una fonte attendibile e quindi il codice iniettato non viene identificato come dannoso. Il contenuto attivo è solitamente uno script (codice di programmazione eseguito dal browser) che può consentire a un utente malintenzionato di accedere a contenuti di pagine sensibili, cookie di sessione o altre informazioni conservate dal browser per conto dell'utente. Con gli attacchi XSS è possibile sottrarre dati sensibili a tutti gli utenti che accedono ai siti infetti. Ad esempio, rubare il SESSIONID di una sessione di un sito web (che si trova nei cookies del browser dell'utente, questi vengono utilizzati per evitare che ogni connessione allo stesso sito richieda un'autenticazione oltre la prima), potrebbe consentire a un impostore di bypassare la fase di login e operare illecitamente per conto dell'utente.

Cross-Site Request Forgery (CSRF)

Un CSRF (Falsificazione di richieste tra siti) è una richiesta http involontaria (ma predisposta dall'attaccante) a un sito Web eseguita da un utente già autenticato su quel sito Web; il sito web eseguirà la richiesta senza ulteriori verifiche, consentendo all'attaccante di portare a termine il suo attacco. Questo tipo di attacco consente di modificare i dati dell'utente o di eseguire transazioni indesiderate, come acquisti indesiderati, se si trattasse ad esempio di un sito di e-commerce. L'attaccante deve avere un'ottima conoscenza del sito in questione per effettuare un attacco CSRF.

Un esempio di attacco CSRF viene eseguito modificando un URL. Ad esempio il codice sorgente della pagina web viene analizzato dal malware per trovare il link "Cambia password" e verificare che venga utilizzato il metodo http GET, quindi è possibile formare un nuovo link con una nuova password (scelta dall'attaccante) e inviarlo alla vittima, ad esempio tramite un allegato e-mail. Se viene cliccato il link e capita che la vittima sia già autenticata sul sito, si attiva una procedura nascosta di avvio della modifica della password, cedendo di fatto l'account all'hacker. I metodi corretti di modifica della password richiedono di inserire anche la vecchia password (e il metodo GET non viene utilizzato). Inoltre, l'uso di un metodo CAPTCHA (un test per verificare se l'utente è un essere umano o un computer) è abbastanza efficace.

Attacchi relativi ai file

Esistono molti tipi di attacchi ai file. Le vulnerabilità più note sono:

Insecure Direct Object References: si verifica quando un'applicazione richiede una risorsa da un server chiamandola con il suo nome, ma consentendo all'utente di modificare quest'ultimo per richiedere un'altra risorsa.

Local File Inclusion (LFI): sfrutta la vulnerabilità delle funzioni GET e POST del linguaggio PHP per inserire file dannosi nel codice.

Path Traversal: consente a un utente malintenzionato di accedere a directory con accesso limitato ed eseguire comandi. Spesso vengono utilizzati in seguito a un LFI per penetrare nel sistema e ottenere i privilegi necessari.

Per proteggersi da questi tipi di attacchi, è importante impedire il caricamento di file sui server, disabilitare l'esecuzione di codice remoto, limitare l'accesso di PHP al file system e aggiornare periodicamente il software.

Denial of Service (DoS) and Distributed Denial of Service (DDoS)

Forse i più noti attacchi di rete, gli attacchi DoS consistono fondamentalmente nell'invio di molti pacchetti per bloccare la possibilità di ricevere altri pacchetti, impedendo così l'accesso alla macchina agli utenti regolari, bloccando i servizi offerti.

La differenza tra un attacco DoS e un attacco DDoS sta nel fatto che nel primo caso viene utilizzato un singolo host per effettuare l'attacco, mentre nel secondo l'attacco viene effettuato con molti host (chiamati zombie) contemporaneamente. Questi host sono stati precedentemente infettati da un malware e al comando dell'attaccante avviano la connessione alla macchina vittima.

Tra gli attacchi DoS, possiamo ricordare:

- **Syn-Flood:** consiste nell'invio di una grande quantità di richieste di connessione. La macchina che riceve una richiesta (un pacchetto SYN) da un client risponde al client inviando un messaggio, allocando alcune risorse per questa connessione e aspettando il pacchetto di conferma (che non arriverà mai). In questo modo la richiesta di connessione rimane attiva finché non si verifica un timeout. Se tutte le connessioni possibili sono state avviate, il servizio fornito dalla macchina non sarà più disponibile. Il timeout avviene in pochi secondi, ma l'enorme quantità di richieste mantiene la macchina intasata.
- **Smurf:** questo attacco è una combinazione di spoofing IP e richieste ICMP. Per spoofing IP si intende l'invio di un pacchetto IP in cui l'IP del mittente è falsificato. Inviando continuamente echo request ICMP (il messaggio inviato dal comando ping, utilizzato per verificare se una macchina è raggiungibile) a indirizzi IP di broadcasting dopo aver sostituito l'indirizzo IP del mittente con quello della vittima, la macchina vittima riceverà una risposta per ogni richiesta da qualsiasi host attivo sulla rete, bloccando così la connessione (flood). Per evitare che questi attacchi si verifichino, è possibile disabilitare nei router l'inoltro degli indirizzi IP di

broadcast.

- **Ping of death:** consiste nell'invio di un pacchetto frammentato con una dimensione IP totale superiore al massimo consentito a una macchina vittima; una volta che la macchina riassume il pacchetto può incorrere in buffer overflow.

È comune che un attacco DDoS utilizzi una botnet, ossia migliaia o milioni di computer infettati da malware che, su comando dell'aggressore, eseguono all'unisono un attacco Syn-Flood su un server vittima.

Man in The Middle (MITM)

Con MITM si intende un'intera tipologia di attacchi che mirano a intercettare, analizzare e ritrasmettere una comunicazione tra due parti che credono di comunicare direttamente tra loro.

Questo tipo di attacco viene solitamente effettuato utilizzando lo spoofing IP: gli sniffer di rete intercettano una connessione e vengono creati nuovi pacchetti cambiando l'indirizzo IP del mittente e del destinatario, in modo che il mittente invii i dati all'attaccante che a sua volta li reinvierà al destinatario e viceversa.

Esistono molti modi per difendersi da questi tipi di attacchi, il più efficace e utilizzato oggi è l'adozione della crittografia end-to-end, possibilmente con certificati digitali.

ARP Spoofing or ARP Poisoning

L'Address Resolution Protocol (ARP) consente alle comunicazioni di rete di raggiungere un dispositivo specifico sulla rete locale. Viene utilizzato per mappare gli indirizzi IP agli indirizzi MAC (Media Access Control) utilizzati per identificare gli host di una LAN. Un host deve utilizzare il protocollo ARP per conoscere l'indirizzo MAC del router/gateway per connettersi a Internet. Ogni host mantiene una cache ARP per mappare le coppie IP/MAC già richieste e risolte. Se l'host non conosce l'indirizzo MAC di un certo indirizzo IP locale, invia un pacchetto di richiesta ARP in broadcast, chiedendo a tutti gli altri host della rete locale chi possiede l'indirizzo IP. L'host con quell'indirizzo IP risponderà, fornendo così il suo indirizzo MAC.

Un attacco ARP spoofing/poisoning può essere classificato come un attacco MITM su rete locale. L'attaccante deve avere accesso alla rete locale e la scansiona per determinare gli indirizzi IP del router e di un altro host. L'attaccante invia risposte ARP contraffatte che fanno credere al router e all'host che l'indirizzo MAC dell'attaccante sia quello dell'host e del router rispettivamente. In questo modo, il router e la workstation si collegano al computer dell'aggressore, anziché l'uno all'altro. Quando i due dispositivi aggiornano le voci della cache ARP, l'host e il router comunicano con l'aggressore invece che direttamente tra loro.

DNS Spoofing

Il Domain Name System (DNS) è il sistema utilizzato per convertire i nomi Internet (più propriamente FQDN - Fully Qualified Distinguished Names, ad esempio `www.server.net`) nel corrispondente indirizzo IP. Il sistema è gerarchico e decentralizzato con alcuni server di top level. Ogni volta che un sistema tenta di raggiungere una risorsa di rete, viene effettuata una connessione a un server DNS per convertire il nome della risorsa in un indirizzo IP. I record delle risorse contengono solitamente informazioni amministrative e altri dati utilizzati, ad esempio, dai server di posta.

L'alterazione dei record DNS permette di deviare il traffico di rete dal sito legittimo a un sito Web falso ("spoofed"). Il server spoofed potrebbe quindi effettuare attacchi MITM, spoofing, ecc. Per prevenire il DNS spoofing i server DNS devono essere aggiornati, ma esistono altri metodi più complessi di firma dei record.

Attacchi Drive-by

In un attacco drive-by, il codice dannoso viene incorporato in un sito Web non sicuro. Quando la pagina infetta viene visitata, lo script viene automaticamente eseguito sul computer della vittima e lo compromette, non è necessario fare nulla sul sito o inserire alcuna informazione. In questo caso l'uso di un antivirus e/o firewall personale aggiornato, recente e decente è l'unica protezione in quanto possono rilevare dal proprio database se un sito non è sicuro prima che un utente lo visiti.

Riutilizzo delle credenziali

La necessità di disporre di credenziali per diversi servizi porta gli utenti a riutilizzare alcune delle loro password. Una volta che un aggressore dispone di un elenco di nomi utente e password di siti web o servizi compromessi (ad esempio ottenuti dal Dark Web), è possibile che la stessa coppia utente/password venga utilizzata dallo stesso utente su altri sistemi. Questo riduce notevolmente il numero di password da controllare con un attacco a dizionario. Per contrastare questo tipo di attacco sono efficaci le stesse contromisure adottate per contrastare gli attacchi alle password. Quando sono necessarie molte password complesse, alcuni trovano utili i gestori di password.

Web Session Hijacking

Un attacco di tipo Session Hijacking (dirottamento) consiste nello sfruttamento dei token di sessione, utilizzati dal browser Web come meccanismi di controllo della sessione. Il server Web deve mettere in relazione diverse connessioni dello stesso utente, poiché per ogni pagina Web viene creata una connessione TCP per ogni elemento (testo, immagini, ecc.). Il metodo più

utilizzato si basa su un token che il server Web invia al browser client dopo che l'autenticazione del cliente ha avuto successo. Un token di sessione è solo un identificatore, una lunga stringa univoca, e può essere utilizzato per evitare la necessità di eseguire una procedura di autenticazione per ogni connessione. Questo attacco viene eseguito rubando o prevedendo un token di sessione valido per ottenere un accesso non autorizzato al server Web.

TCP Session Hijacking

Le connessioni TCP devono garantire che i pacchetti vengano consegnati nel giusto ordine sequenziale; per ottenere questo risultato, il protocollo prevede l'utilizzo di pacchetti di acknowledgement e sequence numbers. Questi pacchetti di acknowledgement sono solo numeri che vengono incrementati in ogni pacchetto in qualche modo dopo l'handshake iniziale a 3 vie. L'obiettivo del dirottare la sessione TCP è creare uno stato in cui il client e il server non sono in grado di scambiare dati, consentendo di produrre falsi pacchetti accettabili per entrambe le parti. In questo modo, l'aggressore è in grado di ottenere il controllo della sessione. Se la sequenza può essere prevista, l'aggressore può inviare pacchetti contraffatti che vengono riconosciuti come validi dall'host vittima.

IP Spoofing

Un IP spoofing viene solitamente eseguito con un attacco DoS: il server legittimo viene bloccato da un numero enorme di connessioni e un server falso imposta il proprio indirizzo IP con lo stesso numero del server legittimo. In questo modo l'utente viene guidato verso il server falso che invia messaggi con un indirizzo IP che indica che il messaggio proviene da un host affidabile.

Nei casi in cui l'instradamento alla fonte è disabilitato, il dirottatore di sessione può anche utilizzare il *blind hijacking*, in cui inietta i suoi dati dannosi nelle comunicazioni intercettate nella sessione TCP. Si chiama blind perché non può vedere la risposta; anche se il dirottatore può inviare i dati o i comandi, in pratica indovina le risposte del client e del server.

Zero-day Exploit

In senso stretto, l'exploit zero-day non è un attacco specifico, in quanto i criminali informatici vengono a conoscenza di una vulnerabilità appena scoperta in determinati software e sistemi operativi e quindi prendono di mira le organizzazioni che utilizzano quel software prima che sia disponibile una patch.

Leak of Code

Il codice sorgente delle applicazioni può talvolta contenere informazioni sensibili e fornire spunti preziosi per un attacco informatico, oltre al valore economico intrinseco del codice stesso.



Co-funded by the
Erasmus+ Programme
of the European Union

Questions:

In file separato.

Attachments:

In file separato

References:

Ogni argomento ha una sua pagina su Wikipedia, a volte già troppo approfondita per gli scopi di questo corso; i libri sono troppo specializzati e sono utili solo dopo una formazione molto approfondita sull'argomento specifico.

1. Malware, or malicious software
URL: <https://www.malwarebytes.com/malware>
2. Cybersecurity for everyone.
<https://www.avast.com/c-malware>
3. What are Cyber Security Threats?
URL: <https://www.imperva.com/learn/application-security/cyber-security-threats/>
4. 21 Top Cybersecurity Threats and How Threat Intelligence Can Help
URL: <https://www.exabeam.com/information-security/cyber-security-threat/>