

Czym jest NIS2?	1
Czym jest SOC	2
Czym jest SIEM.....	2
Definicje	2
Jak działa SIEM	3
Agregacja	3
Przetwarzanie i Normalizacja.....	3
Korelacja	4
Prezentacja	4
Reakcja	4
Rodzaje ataków i zagrożeń	4
Co to jest atak DDos?.....	4
Co to jest Atak Brute Force	5
Co to jest Phishing	6
Co to jest Threat Hunting	8
Co to jest Threat Intelligence	8

Czym jest NIS2 ?

NIS2 (Network and Information Systems Directive 2) to dyrektywa unijna, której głównym celem jest zapewnienie wysokiego poziomu cyberbezpieczeństwa państw członkowskich UE poprzez wprowadzenie nowych obowiązków dla firm z kolejnych, kluczowych sektorów gospodarki.

NIS2 dotyczy podmiotów określanych jako świadczące usługi kluczowe oraz ważne dla gospodarki i społeczeństwa Unii Europejskiej. W praktyce oznacza to, że firmy, które działają w tych sektorach, muszą wdrożyć bardziej rygorystyczne niż dotychczas środki bezpieczeństwa, aby chronić swoje systemy przed cyberzagrożeniami i spełniać nowe wymogi prawne. Przepisy dyrektywy NIS2 wymuszają konieczność nowelizacji polskiej ustawy o Krajowym Systemie Cyberbezpieczeństwa (uKSC), która powinna zacząć być stosowana od 18 października b.r.

Czym jest SOC

Jednym z podstawowych wymagań przewidzianych w NIS2 jest zapewnienie odpowiedniego procesu zarządzania incydentami, którego zadaniem jest wykrywanie incydentów oraz minimalizowanie ich skutków i wpływu na odbiorców. W praktyce oznacza to konieczność wdrożenia mechanizmów i narzędzi, które pozwolą zapobiegać i natychmiast reagować na zaistniałe zdarzenia. Zadania te mogą być realizowane przez zespół świadczący usługi cyberbezpieczeństwa tj. Security Operations Center, który prowadzi monitoring zasobów oraz środowisk IT.

Czym jest SIEM

SIEM, czyli Security Information and Event Management to rodzaj oprogramowania używanego działach bezpieczeństwa IT. System SIEM zapewnia całościowy wgląd w to, co dzieje się w sieci w czasie rzeczywistym i wspiera eliminację tych zagrożeń. SIEM łączy zarządzanie incydentami bezpieczeństwa z zarządzaniem informacjami o monitorowanym środowisku i ma on kluczowe znaczenie w procesie kontroli nad tym, co dzieje się w ich sieci w czasie rzeczywistym. Jest to system złożony z wielu komponentów, służący do monitorowania i analizy, którego celem jest pomoc organizacjom w wykrywaniu zagrożeń i łagodzenie skutków ataków.

Definicje

SIEM łączy kilka różnych obszarów i narzędzi w ramach jednego zintegrowanego systemu:

Log Management (LMS) – narzędzia używane do tradycyjnego zbierania i przechowywania logów

Security Information Management (SIM) – narzędzia lub systemy koncentrujące się na gromadzeniu i zarządzaniu danymi związanymi z bezpieczeństwem z wielu źródeł, takimi jak: zapory, serwery DNS, routery, antywirusy

Security Event Management (SEM) – systemy oparte na proaktywnym monitorowaniu i analizie, w tym wizualizacji danych, korelacji zdarzeń i alarmowaniu

SIEM łączy wszystkie powyższe elementy w jedną platformę, która automatycznie zbiera i przetwarza informacje z rozproszonych źródeł, przechowywać je w jednej scentralizowanej lokalizacji, porównywać różne zdarzenia i generować alerty na podstawie tych informacji.

Jak działa SIEM

SIEM działa poprzez gromadzenie logów i zdarzeń generowanych przez hosty, systemy bezpieczeństwa i aplikacje w całej infrastrukturze organizacji i zestawienie ich na jednej scentralizowanej platformie. Od zdarzeń z oprogramowania antywirusowego do logów z firewall'a, SIEM identyfikuje te dane i dzieli je na kategorie, które na dalszych etapach działań anty zagrożeniowych są bardzo pomocne. Gdy oprogramowanie wykryje aktywność, która może oznaczać zagrożenie dla organizacji, generowane są ostrzeżenia w celu wskazania potencjalnego problemu i szybkie powiadomienie odpowiednich działów bezpieczeństwa.

Alerty można ustawić z niskim lub wysokim priorytetem przy użyciu zestawu wstępnie zdefiniowanych reguł. Na przykład, jeśli konto użytkownika wygeneruje 10 nieudanych prób logowania w ciągu 10 minut, może to zostać oznaczone jako podejrzanе działanie, ale ustawione na niższy priorytet, ponieważ najprawdopodobniej jest to użytkownik, który zapomniał swojego hasła do konta. Jeśli jednak konto doświadczy 100 nieudanych prób zalogowania w przeciągu 5 minut, najprawdopodobniej będzie to atak brute-force i incydent zostanie oznaczony z wysokim priorytetem.

Funkcjonalności SIEM

Agregacja

„Logi” są doskonałym źródłem danych zapewniającym szczegółowy obraz tego, co dzieje się w czasie rzeczywistym. Jest to podstawowe źródło informacji dla SIEM. Niezależnie od tego, czy są to dane z innych systemów bezpieczeństwa, czy błędy z działania usług, SIEM gromadzi je i przechowuje w jednej centralnej lokalizacji. Proces zbierania danych wykonywany jest zwykle przez narzędzia wdrożone na monitorowanych systemach i skonfigurowane do przesyłania danych do wspólnej bazy.

Przetwarzanie i Normalizacja

Największym wyzwaniem w zbieraniu danych w kontekście SIEM jest różnorodność formatów danych. Każde z urządzeń generuje i wysyła dane w swoim zdefiniowanym formacie.

Aby umożliwić wydajną interpretację danych z różnych źródeł, systemy SIEM są w stanie znormalizować dane. Proces ten obejmuje przetwarzanie logów w czytelny i uporządkowany format, wyodrębnienie z nich istotnych informacji oraz mapowanie różnych pól i wartości, które zawierają.

Korelacja

Po zebraniu, przetworzeniu i zapisaniu danych z logów kolejnym krokiem jest korelacja zdarzeń. Sprowadza się to do wyłuskania istotnych informacji z infrastruktury organizacji i powiązanie ich z incydem bezpieczeństwa. Działanie korelacyjne opiera się na regułach wbudowanych system SIEM, predefiniowanych scenariuszach ataków lub tworzonych i konfigurowanych algorytmach. Reguła korelacji określa sekwencję zdarzeń, która może wskazywać na incydent bezpieczeństwa. Na przykład można utworzyć regułę, że jeśli więcej niż 1000 żądań jest wysyłanych z określonych adresów IP i określonych portów w określonym przedziale czasu, to jest to związane z atakiem DDoS.

Prezentacja

Możliwość wizualizacji danych i zdarzeń jest kolejną kluczową funkcjonalnością systemów SIEM. Pomaga zauważać trendy czy anomalie i monitorować ogólny stan środowiska IT.

Reakcja

Po właściwej korelacji i monitorowaniu zdarzeń, aby zapewnić kompleksową ochronę systemu należy mieć sposób postępowania z incydentami po ich wykryciu. Służą temu mechanizmy automatycznego blokowania i ograniczania dostępu do zaatakowanego urządzenia czy zasobu. Możliwe jest też automatyczne wykonywanie akcji takich jak wywołanie skryptu, utworzenie zgłoszenia serwisowego czy wysłanie maila.

Rodzaje Ataków i Zagrożeń

Co to jest atak DDos ?

Ataki typu DDoS (ang. distributed denial of service, w wolnym tłumaczeniu: rozproszona odmowa usługi) są jednymi z najczęściej występujących ataków hakerskich, które kierowane są na systemy komputerowe lub usługi sieciowe i mają za zadanie zajęcie wszystkich dostępnych i wolnych zasobów w celu uniemożliwienia funkcjonowania całej usługi w sieci Internet (np. Twojej strony internetowej i poczty znajdującej się na hostingu).

Na czym polega atak DDoS?

Atak DDoS polega na przeprowadzeniu ataku równocześnie z wielu miejsc jednocześnie (z wielu komputerów). Atak taki przeprowadzany jest głównie z komputerów, nad którymi przejęta została kontrola przy użyciu specjalnego oprogramowania (np. boty i trojany). Oznacza to, że właściciele tych komputerów mogą nawet nie wiedzieć, że ich komputer, laptop lub inne urządzenie podłączone do sieci, może być właśnie wykorzystywane, bez ich świadomości, do przeprowadzenia ataku DDoS. Najczęściej przejmowanymi komputerami w celu wykorzystania ich do ataku DDoS, są komputery, smartfony oraz inne urządzenia podłączone do Internetu, które nie zostały zabezpieczone programem antywirusowym. Pamiętaj, aby zabezpieczyć swój sprzęt programem antywirusowym,

możesz go zamówić w home.pl (np. Norton Security, Bitdefender, Avast Internet Security). Sprawdź również propozycje antywirusów na telefon. Atak DDoS rozpoczyna się w momencie, gdy wszystkie przejęte komputery zaczynają jednocześnie atakować usługę WWW lub system ofiary. Cel ataku DDoS zasypywany jest wtedy fałszywymi próbami skorzystania z usług (np. mogą być to próby wywołania strony internetowej lub inne żądania). Najczęściej w następstwie ataku DDoS, usługi pocztowe lub serwisy internetowe mogą nie być widoczne w sieci Internet.

Dlaczego atak DDoS powoduje przerwy w funkcjonowaniu usług?

Każda próba skorzystania z usługi (np. próba wywołania strony internetowej) wymaga, aby atakowany komputer przydzielił odpowiednie zasoby do obsługi tego żądania (np. procesor, pamięć, pasmo sieciowe), co przy bardzo dużej liczbie takich żądań, prowadzi do wyczerpania dostępnych zasobów, a w efekcie do przerwy w działaniu lub nawet zawieszenia atakowanego systemu.

Jak chronić się przed atakami DDoS?

Ataki typu DDoS są obecnie najbardziej prawdopodobnym zagrożeniem dla firm działających w sieci, a ich konsekwencje sięgają dalej niż tylko obszaru IT, ale również powodują realne, mierzalne straty finansowe i wizerunkowe. Ataki tego typu ciągle ewoluują i stają się coraz bardziej precyzyjne. Ich celem jest zużycie wszystkich dostępnych zasobów infrastruktury sieciowej lub łącza internetowego. Ochrona przed atakami DDoS odbywa się poprzez zmianę ustawień DNS, co powoduje skierowanie całego ruchu HTTP/HTTPS przez warstwę filtrującą, w której to dokonywana jest szczegółowa inspekcja każdego pakietu i zapytania.

C o to jest Atak Brute Force

Brute Force to sposób na uzyskanie danych dostępowych, znalezienie ukrytych stron internetowych (niepublicznych, wyłączonych z indeksowania) czy też odkrycie klucza używanego do szyfrowania danych. Atak Brute Force oparty jest na metodzie prób i błędów, których celem jest odgadnięcie danych. Ta metoda ataku na dane użytkownika jest bardzo stara i niestety prosta do wykonania.

Brute Force opiera się na metodzie zgadywania, co może zająć bardzo dużo czasu, biorąc pod uwagę ilość kombinacji loginów, haseł, adresów, kluczy szyfrujących, nie wspominając o ich nieznanej długości. Hakerzy opracowali liczne narzędzia do szybszego przeprowadzania takich ataków.

Brute Force i automatyzacja procesu

W działaniu ataków Brute Force pomagają oprogramowanie, które wykorzystując dostępną moc procesorów oraz procesorów graficznych, przyspiesza cały proces. Ten sam w sobie może mimo wszystko trwać kilka dni, tygodni, a nawet miesięcy, w zależności od skomplikowania klucza szyfrującego. Oprogramowanie generuje dużą liczbę kolejnych kombinacji w oparciu o liczne techniki i algorytmy, np. dopasowywania

słów, ciągów znaków, zamiany liczb na litery, tworzenia popularnych skojarzeń, itp.. Brute Force zgaduje dane dostępu, tworząc nieskończenie wiele kombinacji.

Metoda słownikowa w ataku brute force

Ataki tego typu można przeprowadzić, korzystając na przykład z tak zwanej metody słownikowej, a więc dysponując, np. loginem lub hasłem, poszukiwać drugiej zmiennej na podstawie popularnych fraz. Metoda ta jest o tyle skuteczna, o ile użytkownicy nadal używają popularnych słów, imion, nazw jako hasła dostępu. W metodzie słownikowej najczęściej znamy jedną zmienną, np. login, który został pozyskany w wyniku wycieku danych. Następnie metodą prób i błędów, zaczynając przykładowo od ww. metody słownikowej, rozpoczyna się zgadywanie hasła dostępu do konta, panelu, usługi, danych, itp.. W drugą stronę, dysponując bazą haseł, popularnością występowania fraz, poszukiwane są loginy.

Jak rozpoznać atak brute force?

Atak Brute Force jest atakiem, który przesyła bardzo wiele danych na serwer, na którym ma dojść do odszyfrowania danych. Z tego powodu nim nastąpi samo złamanie zabezpieczeń, można zaobserwować, np. wolniejsze działanie strony WWW lub innego mechanizmu logowania. Niewykluczone, że unieruchomiony zostanie cały serwer. Atak typu Brute Force można przewidzieć i zatrzymać dzięki bieżącemu monitoringowi: obserwując m.in. anomalie w ruchu na serwerze, np. w działaniu mechanizmów logowania. Sygnały, które mogą wskazywać na atak brute force, to m.in.:

- wielokrotne nieudane próby logowania z tego samego adresu IP na wiele różnych kont,
- logowanie do konta z nietypowego dla niego adresu IP,
- inne niż przeważnie zachowanie użytkownika po udanym logowaniu.

Często nie zdajemy sobie sprawy z tego, że początek większego „ataku” może tkwić w naszej aktywności w internecie oraz poza nim. Ataki brute force nie muszą być przeprowadzane wyłącznie tam, gdzie liczy się zdobycie informacji z Twojego konta. Ataki takie mogą stanowić ciąg prac, zmierzających do uzyskania cennej informacji. Złamane hasło dostępu do systemu Windows, przejęcie konta pocztowego, identyfikacja adresu e-mail z systemem płatności, dostęp do kart kredytowych, nieautoryzowane zakupy.

Co to jest Phishing

Phishing to rodzaj cyberataku, podczas którego cyberprzestępca próbuje wyłudzić od ofiary poufne informacje. Jest to najprostszy i jednocześnie najpopularniejszy rodzaj cyberataku, którego celem jest np. kradzież loginów, numerów kart kredytowych i rachunków bankowych czy wrażliwych informacji firmowych. Podczas ataku tego typu, haker może również zainfekować komputer użytkownika oprogramowaniem typu malware. Celem phishingu nie jest urządzenie, ale człowiek i posiadane przez niego dane. Cyberprzestępca bardzo często podszywa się pod zaufaną instytucję lub osobę zaufania społecznego, autorytet dla ofiary, wymuszając strach czy pośpiech, aby skłonić ją do szybkiego działania.

Przykładem takiej wymuszonej natychmiastowej reakcji może być logowanie do konta bankowego w celu autoryzacji danych.

Ataki phishingowe możemy podzielić na cztery rodzaje:

1. Cyberprzestępca wysyła w wiadomości e-mail link, licząc na to, że ofiara w niego kliknie (phishing e-mail).
2. Haker próbuje wciągnąć ofiarę w rozmowę telefoniczną (vishing).
3. Haker wysyła wiadomość SMS z prośbą o kliknięcie w odnośnik lub oddzwonienie do nadawcy (smishing).
4. Spear phishing (oraz jego odmiana – whaling),

Phishing – przykłady

Phishing może polegać m.in. na jednym z poniższych działań:

- Wysyłaniu wiadomości z zainfekowanymi załącznikami np. do klientów banków. Kliknięcie w taki załącznik może skutkować pobraniem złośliwego oprogramowania, które szyfruje dane.
- Tworzeniu fałszywych stron www (np. strona logowania do banku) podobnych do oryginałów, na której użytkownik wpisuje login i hasło.
- Wysyłaniu wiadomości SMS, które informują o konieczności uiszczenia jakiejś opłaty, np. za zamówioną rzekomo przesyłkę.
- Rozsyłaniu fałszywych wezwań do zapłaty z groźbą o wszczęciu postępowania windykacyjnego.
- Wyłudzeniu kodu BLIK przez osoby, które podszywają się pod znajomych z portali społecznościowych ofiary (np. z Facebooka) itp.

Co to jest oprogramowanie Ransomware

Ransomware to rodzaj złośliwego oprogramowania, które szyfruje ważne pliki przechowywane na dysku lokalnym i sieciowym oraz żąda okupu za ich rozszyfrowanie. Hakerzy tworzą tego typu oprogramowanie w celu wyłudzenia pieniędzy prze z szantaż.

Oprogramowanie ransomware jest zaszyfrowane, co oznacza, że nie da się zdobyć klucza, a jedynym sposobem na odzyskanie informacji jest przywrócenie ich z kopii zapasowej.

Sposób działania oprogramowania ransomware sprawia, że jest ono wyjątkowo szkodliwe. Inne rodzaje malware niszczą lub kradną dane, ale nie zamykają drogi do ich odzyskania. Natomiast w przypadku ransomware, jeśli zaatakowany podmiot nie ma kopii zapasowej danych, aby je odzyskać, musi zapłacić okup. Zdarza się, że firma płaci okup, a haker i tak nie przekazuje klucza do rozszyfrowania.

Program typu ransomware rozpoczyna działanie od przeskanowania lokalnych i sieciowych magazynów danych w poszukiwaniu plików do zaszyfrowania. Na cel bierze pliki, które uzna za ważne dla firmy lub osoby prywatnej. Zaliczają się do nich także pliki

kopii zapasowych, które mogłyby posłużyć do odzyskania informacji. Poniżej znajduje się lista niektórych typów plików poszukiwanych przez ransomware:

- Microsoft Office: .xlsx, .docx, i .pptx oraz starsze wersje
- Obrazy: .jpeg, .png, .jpg, .gif
- Rysunki techniczne: .dwg
- Dane: .sql i .ai
- Wideo: .avi, .m4a, .mp4

Różne rodzaje oprogramowania ransomware biorą na cel różne typy plików, ale istnieje pewien zbiór wspólny. Większość programów ransomware szuka plików Microsoft Office, ponieważ firmy często przechowują w nich najważniejsze informacje. Zaszzyfrowanie ważnych plików zwiększa szansę na uzyskanie okupu.

Co to jest Threat Hunting

Raporty wskazują, iż cyberprzestępcy często czają się tygodnie, a nawet miesiące w sieci organizacji, zanim zostaną wykryci lub sami się ujawnią. Cierpliwie czekają, aby wykraść informacje lub zdobyć dane uwierzytelniające do odblokowania dalszych dostępów. Threat hunting jest to aktywne wyszukiwanie intruzów „zagnieżdżonych”

w infrastrukturze IT. Są to proaktywne działania „śledcze”, zmierzające do wykrycia zidentyfikowania intruzów wewnętrznych i zewnętrznych. Podstawowym założeniem jest, iż doszło do naruszenia zasad bezpieczeństwa, a intruz jest obecny i ukrywa się w sieci firmowej. Niewykrycie na czas atakującego umożliwia mu utrzymanie dostępu do systemów kontrolnych, zwiększenie uprawnień i uodpornienie na skuteczną detekcję i eliminację.

Co to jest Threat Intelligence

Threat Intelligence to stałe pozyskiwanie informacji z zewnętrznych źródeł o zagrożeniach dla firm. Platforma dostarcza informacje wykorzystywane do automatycznego wzbogacania wewnętrznych systemów monitoringu typu SIEM. Najprostszym przykładem takiego wzbogacania może być pobieranie informacji o adresach IP wykorzystywanych przez atakujących lub też wykrywanie zmian w otwartych portach infrastruktury organizacji.

Jest to zbiór danych gromadzonych, przetwarzanych oraz analizowanych w celu zrozumienia motywów, celów i zachowań cyberprzestępców. Taka baza umożliwia podejmowanie szybszych i bardziej trafnych decyzji dotyczących bezpieczeństwa oraz zmianę kierunków działań i nadanie im charakteru proaktywnego zapobiegania zagrożeniom.