

ISSN 2300-5149

PRZEGLĄD TELEINFORMATYCZNY

W. Kwiatkowski	
Wykrywanie anomalii bazujące na wskazanych przykładach	3
Ł. Strzelecki, A. Ludew	
Koncepcja zwiększenia bezpieczeństwa transakcyjnych aplikacji internetowych	23
K. Liderman, A.E. Patkowski	
Zintegrowane badanie i ocena stanu ochrony zasobów informacji	39
M. Marczyk	
Cyberprzestrzeń jako nowy wymiar aktywności człowieka – analiza pojęciowa obszaru	59
Informacje dla autorów – procedura recenzowania artykułu i sposób przygotowania tekstu dla redakcji Przeglądu Teleinformatycznego	73

PRZEGLĄD TELEINFORMATYCZNY
TELEINFORMATICS REVIEW

Dawniej: BIULETYN INSTYTUTU AUTOMATYKI I ROBOTYKI WAT
(ISSN 1427-3578)
Ukazuje się od 1995 r.

RADA NAUKOWA

Lt. Col. Janos Balogh MSc
dr hab. inż. Antoni M. Donigiewicz – przewodniczący
Hacene Fouchal, PhD
prof. Lech J. Janczewski, DEng
prof. dr hab. inż. Włodzimierz Kwiatkowski
prof. dr hab. inż. Bohdan Macukow
Lt. Col. Lajos Mucha PhD
prof. ing. Vladimír Olej, CSc.



Ministerstwo Nauki
i Szkolnictwa Wyższego

UMOWA NR 630/P-DUN/2018 Z MNiSW (z dn. 21.06.2018 r.) na lata 2018-2019:

- wzrost liczby artykułów anglojęzycznych;
- obsługa strony internetowej czasopisma na platformie IC, nadawanie artykułom naukowym numerów DOI i wykorzystanie profesjonalnego panelu edycyjnego.

Wymienione zadania finansowane są w ramach umowy nr 630/P-DUN/2018 ze środków Ministra Nauki i Szkolnictwa Wyższego przeznaczonych na działalność upowszechniającą naukę.

ADRES REDAKCJI

Redakcja Przeglądu Teleinformatycznego
00-908 Warszawa 46, ul. Gen. W. Urbanowicza 2
tel. 261 83 95 52, fax. 261 83 71 44
e-mail: pt [at] ita.wat.edu.pl
WWW: <http://przeglad.ita.wat.edu.pl/>

Wersją pierwotną czasopisma jest wersja elektroniczna

REDAKTOR NACZELNY:

Antoni Donigiewicz

REDAKTOR WYDANIA

Antoni Donigiewicz

OPRACOWANIE STYLISTYCZNE

Renata Borkowska

PROJEKT OKŁADKI

Barbara Chruszczyk

WYDAWCA: Instytut Teleinformatyki i Automatyki WAT

ISSN 2300-5149

ISSN 2353-9836 (on-line)

Wykrywanie anomalii bazujące na wskazanych przykładach

Włodzimierz KWIATKOWSKI

Instytut Teleinformatyki i Automatyki, Wydział Cybernetyki, WAT,
ul. Gen. W. Urbanowicza 2, 00-908 Warszawa
wlodzimierz.kwiatkowski@wat.edu.pl

STRESZCZENIE: W artykule rozpatrywany jest problem wykrywania anomalii na podstawie zarejestrowanych obserwacji zachowania systemu. Problem jest sformułowany jako zadanie rozpoznawania wzorców zachowania normalnego i zachowania nietypowego. Obydwa wzorce są określone przez wskazanie odpowiednich przykładów. Osobliwość rozwiązywanego zadania wynika z faktu, że zwykle liczebność przykładów jest dużo mniejsza od wymiaru wektora obserwacji. W artykule zostały przedstawione dwie metody detekcji anomalii bazujące na wyznaczaniu rzutów obserwacji na podprzestrzenie wzorców. Wyróżnikiem pierwszej metody jest wykorzystywanie odległości wektora obserwacji od podprzestrzeni wzorców. Druga metoda polega na przeniesieniu zadania rozpoznawania wzorców do podprzestrzeni wzorców.

SŁOWA KLUCZOWE: wykrywanie anomalii, rozpoznawanie wzorców, eksploracja danych, odległość Mahalanobisa

1. Wprowadzenie

Wykrywanie anomalii (nieprawidłowości) należy do podstawowych problemów administrowania systemami, w tym komputerowymi i teleinformatycznymi. Jest to także zasadniczy problem szeroko rozumianej diagnostyki (m.in. technicznej, medycznej). Wykrywanie anomalii nie jest jednak tożsame z testowaniem systemu. Celem testowania systemu jest wykrywanie błędów w systemie i sprawdzenie jego wymaganej funkcjonalności. Przykładem może być testowanie (weryfikacja, walidacja) oprogramowania. Testowanie systemu ma zwykle charakter interaktywny. Polega to na podawaniu do systemu określonych wymuszeń i porównywaniu uzyskiwanej reakcji systemu z reakcją oczekiwaną, zwykle wcześniej określoną w dokumentacji systemu. Wykrywanie

anomalii ma charakter pasywny i jej celem jest sprawdzenie, czy zachowanie systemu nie zmieniło się w stosunku do jego normalnej pracy.

Najczęściej wykrycie anomalii sprowadza się do stwierdzenia niezgodności obserwacji z modelem (regułami) działania systemu. Takie podejście obejmuje więc przypadki zaobserwowania wartości odstających (*outliers*, *discordant observations*), wyjątków (*exceptions*), osobliwości (*peculiarities*) czy też po prostu nowych zachowań (*novelty items*). Konsekwencją stwierdzenia anomalii może być konieczność odstąpienia od zarządzania (sterowania) bazującego na wykorzystywanym modelu (i np. przejście do procedury obsługi wyjątków). Zaobserwowanie wartości odstających zwykle skutkuje nieuwzględnianiem ich w formułowaniu modelu działania systemu. Pozostałe przypadki wymagają najczęściej modyfikacji lub zmiany modelu.

Wykrywanie anomalii wymaga uprzedniego określenia, jakie obserwacje (pomiar, cechy) będą podstawą wnioskowania. Od trafności przyjętych ustaleń zależy użyteczność orzeczeń.

Wykrycie anomalii z reguły następuje w wyniku weryfikacji hipotezy „zachowanie jest normalne”. Takie podejście wymaga zdefiniowania specjalnego modelu normalności. Model ten jest rozumiany jako zestaw reguł do orzeczenia zgodności obserwacji z przyjętym modelem działania systemu. W szczególnym przypadku budowa takiego modelu może polegać na wyznaczeniu obszaru normalności w przestrzeni obserwacji (pomiarów, cech).

Weryfikacja hipotezy „zachowanie jest normalne” przeciwko alternatywnej hipotezie złożonej „zachowanie nie jest normalne” jest zdecydowanie trudniejsza niż w przypadku, gdy hipoteza alternatywna jest prosta (np. „zachowanie nie jest normalne z powodu awarii podzespołu A”). W drugim przypadku problem zostaje zawężony do wykrywania interesującej nas anomalii (nieprawidłowości). Nie można jednak oczekiwać, że uda się skatalogować wszystkie przypadki anomalii i opracować dla nich odpowiednie alternatywne hipotezy proste. Poszukuje się więc rozwiązań kompromisowych między wnioskowaniem na podstawie znanego modelu normalności bez żadnej wiedzy o anomaliiach a wnioskowaniem na podstawie znanego modelu normalności i znanego modelu anomalii. Kompromis ten można osiągnąć, wskazując zaobserwowane przykłady (przypadki) zachowań zarówno normalnych, jak i anormalnych (*anomalous*).

Przegląd metod rozwiązywania problemu wykrywania anomalii jest przedstawiony w artykule [5]. Przyjęte w pracy [2] sformułowanie problemu wykrywania anomalii obejmuje jako przypadki szczególne wykrywanie nowych zachowań (*novelty detection*) [6] oraz wykrywanie zachowań odstających (*outliers*) [3]. W tych szczególnych przypadkach wymagane jest określenie jedynie wzorca normalności, a zachowanie z nim niezgodne można określić jako anormalne (*abnormal*). Problem wykrywania zachowań anormalnych pojawia

się w przypadku analizy obserwacji nieoznakowanych (*unsupervised outlier detection*) [1, 9].

2. Wykrywanie anomalii na podstawie przykładów

W niniejszym artykule zadanie wykrywania anomalii jest formułowane jako zadanie rozpoznawania wzorców: normalności i anomalii. Sformułowanie tego zadania jest z natury rzeczy trudne. Model statystyczny zachowania normalnego może okazać się zbyt ogólny i z tego powodu wnioskowanie bazujące na ocenie zgodności z nim obserwacji (danych) może być zawodne [2].

Podstawowa trudność budowy modelu anomalii wynika z faktu, że nie jest z góry wiadome, jaki rodzaj obserwacji (danych) może ujawnić odstępstwa od stosowanego modelu. Z tego powodu do wykrywania anomalii próbuje się wykorzystywać wszystkie dostępne dane, choćby tylko potencjalnie użyteczne. Takie podejście prowadzi do konieczności analizy zbiorów danych charakteryzujących się dużymi rozmiarami i dużą różnorodnością. Dane tego typu są uzyskiwane przy automatycznym dokumentowaniu różnorodnych działań (*digital footprint*, *digital shadow*). Model generowania tych danych zwykle nie jest znany.

W tej sytuacji stosowanie analizy typu potwierdzającego (*confirmatory data analysis*) i bazującej na modelach statystycznych jest kontestowane [2]. Alternatywą są badania o charakterze wydobywczym, powszechnie określane jako eksploracja danych (*exploratory data analysis*) [7]. Rutynową techniką wykorzystywaną w eksploracyjnej analizie danych do selekcji wyników obserwacji jest analiza głównych składowych (*principal component analysis* – PCA). W szczególności umożliwia ona redukcję liczby współrzędnych wektora obserwacji.

Rozpatrywany w niniejszym artykule problem dotyczy typowej w zadaniach wykrywania anomalii sytuacji, gdy liczba wskazanych przykładów wzorców (zwłaszcza wzorców anomalii) jest mała względem liczby współrzędnych wektora wyników obserwacji. Wtedy przy analizie obserwacji wzorcowych przypadków występują składowe główne, które charakteryzują się zerową wariancją. Z formalnego punktu widzenia celem przedstawianych dalej badań jest uzyskanie metod minimalnoodległościowych rozpoznawania wzorców wtedy, gdy wymiar przestrzeni obserwacji jest większy od wymiaru podprzestrzeni generowanej przez składowe główne o niezerowej wariancji.

3. Przestrzeń wyników obserwacji

Podstawą badań są wyniki obserwacji zachowań analizowanej klasy systemów. Każda obserwacja odnosi się do jednego przykładu zachowania sytemu. Wynikiem każdej obserwacji jest wektor złożony z ustalonej liczby współrzędnych (interpretowanych jako cechy zachowania systemu). Omawiane wyniki są zestawiane w postaci następującej macierzy:

$$\mathbf{W} = [\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_N] \quad (1)$$

gdzie:

$$\mathbf{w}_k = \begin{bmatrix} w_{1,k} \\ w_{2,k} \\ \dots \\ w_{L,k} \end{bmatrix} \quad (2)$$

Parametr N oznacza liczbę wszystkich przykładów zachowań, a parametr L – liczbę współrzędnych wektora obserwacji. Oznaczmy dalej macierz kowariancji obserwacji następująco:

$$\mathbf{R} = \frac{1}{N-1} \sum_{k=1}^N (\mathbf{w}_k - \bar{\mathbf{w}})(\mathbf{w}_k - \bar{\mathbf{w}})^T \quad (3)$$

gdzie:

$$\bar{\mathbf{w}} = \frac{1}{N} \sum_{k=1}^N \mathbf{w}_k \quad (4)$$

Przyjmujemy dalej, że

$$\text{rank}(\mathbf{R}) = L \quad (5)$$

Odległość pomiędzy wektorami $\mathbf{x}$, $\mathbf{y}$ przestrzeni cech R^L będziemy wyznaczać w sposób uwzględniający wielkość rozrzutu (rozproszenia) współrzędnych pomiaru oraz ich wzajemną korelację. Wymaganie to spełnia odległość Mahalanobisa określona wzorem:

$$d(\mathbf{x}, \mathbf{y}) = \sqrt{(\mathbf{x} - \mathbf{y})^T \mathbf{R}^{-1} (\mathbf{x} - \mathbf{y})}, \quad \mathbf{x}, \mathbf{y} \in R^L \quad (6)$$

4. Rozpoznawanie wzorców metodą minimalnej odległości

Wskazania przykładów obserwacji zakwalifikowanych do wzorca o indeksie $h \in \{1, 2, \dots, H\}$ (gdzie: H – liczba wzorców) będziemy dokonywać przez podanie odpowiedniego zbioru indeksów W_h . Rozpatrywany wzorec będzie więc reprezentowany przez następujący zbiór punktów (klastr) w przestrzeni obserwacji:

$$C(W_h) = \{\mathbf{w}_k : k \in W_h\} \quad (7)$$

składający się ze wskazanych przykładów obserwacji. Liczbę elementów tak rozumianego wzorca W_h oznaczmy jako $N_h = \|C(W_h)\|$. Wnioskowanie o podobieństwie obserwacji $\mathbf{x}$ do wzorca W_h bazuje na określeniu odległości punktu $\mathbf{x}$ od klastra $C(W_h)$. Przykładowo, wybierając metodę centroidalną wyznaczania odległości między klastrami, otrzymujemy zależność:

$$D(\mathbf{x}, C(W_h)) = d(\mathbf{x}, \bar{\mathbf{w}}_h) = \sqrt{(\mathbf{x} - \bar{\mathbf{w}}_h)^T \mathbf{R}^{-1} (\mathbf{x} - \bar{\mathbf{w}}_h)} \quad (8)$$

gdzie:

$$\bar{\mathbf{w}}_h = \frac{1}{N_h} \sum_{j \in W_h} \mathbf{w}_j \quad (9)$$

Z uwagi na sposób wyznaczenia macierzy kowariancji $\mathbf{R}$, taka metoda wnioskowania znajduje uzasadnienie tylko wtedy, gdy pomiary odpowiadające wszystkim wzorcom są jednorodne w następującym sensie: odpowiednie klastry różnią się jedynie wartościami oczekiwanymi (a odpowiadające im macierze kowariancji są jednakowe).

W wielu zagadnieniach, a zwłaszcza w przypadku badania anomalii, macierze kowariancji wzorców różnią się. Rozpatruje się wtedy możliwość zróżnicowania sposobu pomiaru odległości stosownie do macierzy kowariancji poszczególnych wzorców [4].

Macierz kowariancji wyznaczoną na podstawie przykładów wzorca W_h oznaczmy następująco:

$$\mathbf{R}_h = \frac{1}{N_h - 1} \sum_{j \in W_h} (\mathbf{w}_j - \bar{\mathbf{w}}_h)(\mathbf{w}_j - \bar{\mathbf{w}}_h)^T \quad (10)$$

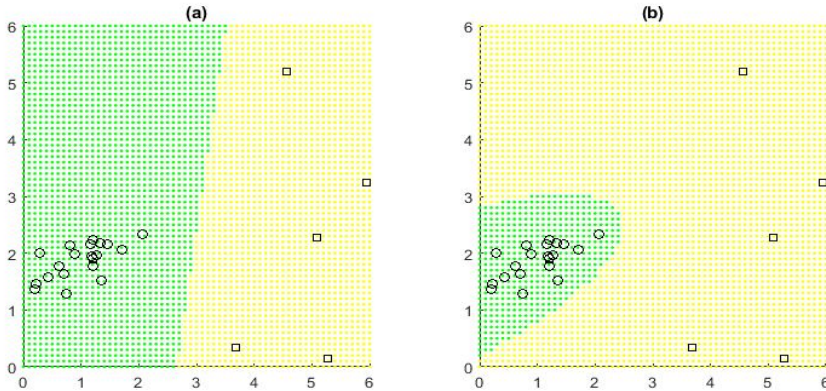
Odległość pomiędzy wektorami $\mathbf{x}$, $\mathbf{y}$ przestrzeni cech R^L zadaną wzorem:

$$d_h(\mathbf{x}, \mathbf{y}) = \sqrt{(\mathbf{x} - \mathbf{y})^T \mathbf{R}_h^{-1} (\mathbf{x} - \mathbf{y})}, \quad \mathbf{x}, \mathbf{y} \in R^L \quad (11)$$

nazywać będziemy dopasowaną do wzorca W_h . Podobnie nazywać będziemy odległość między obserwacją $\mathbf{x}$ a klastrem $C(W_h)$. Przykładowo dla metody centroidalnej odległość ta jest określona wzorem:

$$D_h(\mathbf{x}, C(W_h)) = d_h(\mathbf{x}, \bar{\mathbf{w}}_h) = \sqrt{(\mathbf{x} - \bar{\mathbf{w}}_h)^T \mathbf{R}_h^{-1} (\mathbf{x} - \bar{\mathbf{w}}_h)} \quad (12)$$

Na rysunku 1 przedstawiony jest przykład ilustrujący różnice powodowane wykorzystywaniem odległości dopasowanych do poszczególnych wzorców: W_1 bądź W_2 . Punkty przykładów wzorcowych $C(W_1)$ są przedstawione na rysunku jako kółka, punkty przykładów wzorcowych $C(W_2)$ – jako kwadraty. Punkty przestrzeni obserwacji leżące bliżej przykładów wzorcowych $C(W_1)$ są oznaczone kolorem ciemniejszym. Przedstawiona na rys. 1(a) metoda klasyfikacji prowadzi do uzyskania wyników analogicznych do otrzymywanych w liniowej analizie dyskryminacyjnej (LDA – *linear discriminant analysis*). Rozwiązanie przedstawione na rys. 1(b) ma bezpośrednie odniesienie do kwadratowej analizy dyskryminacyjnej (QDA – *quadratic discriminant analysis*).



Rys. 1.

- (a) Do wyznaczania odległości wykorzystywane są metryki bazujące na macierzy kowariancji obliczonej dla wszystkich obserwacji razem, zgodnie ze wzorem (6).
- (b) Do wyznaczania odległości wykorzystywane są metryki bazujące na macierzy kowariancji obliczanych dla obserwacji każdego wzorca osobno, zgodnie ze wzorem (11)

5. Podprzestrzeń wzorca W_h

Jeśli macierz $\mathbf{R}_h$ jest osobiłwa, obliczenie odległości dopasowanej (12) jest niemożliwe. W takim przypadku proponujemy zredukować liczbę współrzędnych pomiarów tak, aby w uzyskanej w ten sposób podprzestrzeni odpowiednia macierz kowariancji była nieosobiłwa. Podprzestrzeń uzyskaną w ten sposób nazywać będziemy podprzestrzenią wzorca W_h .

Proponujemy dalej, aby redukcję wymiaru wektora pomiaru przeprowadzić w przestrzeni wartości transformat Karhunen–Loève’a [4]. Podstawą przekształcenia Karhunen–Loève’a są ortonormalne wektory własne $\mathbf{t}_k(\mathbf{R}_h)$ macierzy kowariancji $\mathbf{R}_h$. Wektory te spełniają następującą zależność:

$$\mathbf{R}_h \mathbf{t}_k(\mathbf{R}_h) = \lambda_k(\mathbf{R}_h) \mathbf{t}_k(\mathbf{R}_h), \quad k = 1, 2, \dots, L \quad (13)$$

gdzie: $\mathbf{t}_k(\mathbf{R}_h) = \begin{bmatrix} t_{k,1} \\ t_{k,2} \\ \dots \\ t_{k,L} \end{bmatrix}$ a $\lambda_k(\mathbf{R}_h)$ – wartości własne macierzy kowariancji $\mathbf{R}_h$.

Wartości własne $\lambda_k(\mathbf{R}_h)$ są liczbami rzeczywistymi; przyjmujemy, że wartości te są uporządkowane malejąco względem indeksu k . Wtedy macierz przekształcenia Karhunen–Loève’a można przedstawić następująco:

$$\mathbf{T}_h = \begin{bmatrix} \mathbf{t}_1^T(\mathbf{R}_h) \\ \mathbf{t}_2^T(\mathbf{R}_h) \\ \dots \\ \mathbf{t}_L^T(\mathbf{R}_h) \end{bmatrix} \quad (14)$$

Transformaty różnic $\mathbf{w}_k - \bar{\mathbf{w}}_h$ oznaczmy następująco:

$$\mathbf{v}_k = \mathbf{T}_h(\mathbf{w}_k - \bar{\mathbf{w}}_h) \quad (15)$$

gdzie $\bar{\mathbf{w}}_h$ jest określone wzorem (9). Macierz kowariancji $\mathbf{V}_h$ wektorów $\{\mathbf{v}_k, k \in W_h\}$ jest macierzą diagonalną:

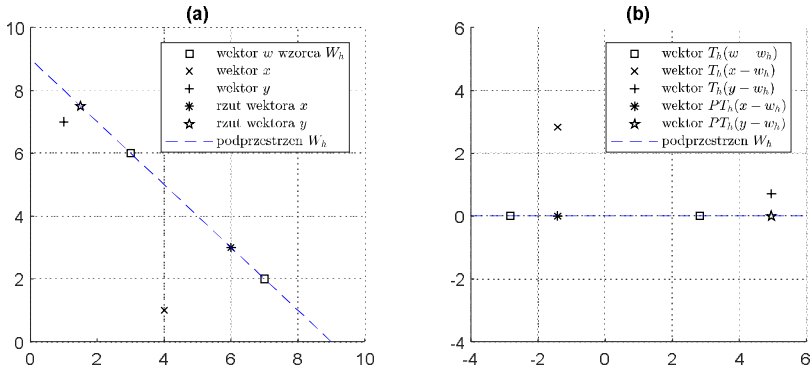
$$\mathbf{V}_h = \frac{1}{N_h - 1} \sum_{k \in W_h} (\mathbf{v}_k - \bar{\mathbf{v}}_h)(\mathbf{v}_k - \bar{\mathbf{v}}_h)^T = \text{diag}(\lambda_1(\mathbf{R}_h), \lambda_2(\mathbf{R}_h), \dots, \lambda_L(\mathbf{R}_h)) \quad (16)$$

przy czym: $\bar{\mathbf{v}}_h = \frac{1}{N_h} \sum_{k \in W_h} \mathbf{v}_k = \frac{1}{N_h} \sum_{k \in W_h} \mathbf{T}_h(\mathbf{w}_k - \bar{\mathbf{w}}_h) = \mathbf{0}$.

Niech $M_h \leq \min\{N_h, L\}$ oznacza liczbę dodatnich wartości własnych macierzy kowariancji $\mathbf{R}_h$. Niech operator $\mathbf{P}$ przekształca wektory $\mathbf{v}_k \in R^L$ w następujący sposób:

$$\mathbf{v}_k = \mathbf{P}\mathbf{v}_k = \begin{bmatrix} 1 & \dots & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots \\ 0 & \dots & 1 & \dots & 0 \\ 0 & \dots & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots \\ 0 & \dots & 0 & \dots & 0 \end{bmatrix} \begin{bmatrix} v_{k,1} \\ \dots \\ v_{k,M_h} \\ v_{k,M_h+1} \\ \dots \\ v_{k,L} \end{bmatrix} = \begin{bmatrix} v_{k,1} \\ \dots \\ v_{k,M_h} \\ 0 \\ \dots \\ 0 \end{bmatrix} \quad (17)$$

Złożenie $\mathbf{PT}_h$ pozwala wyznaczyć poszukiwaną podprzestrzeń wzorca W_h . Ilustracja odwzorowania wektorów przestrzeni cech R^L w wektory podprzestrzeni transformat została przedstawiona na rysunku 2 (w rozpatrywanym przykładzie $L = N_h = 2$, $M_h = 1$). Wektory obserwacji $\mathbf{x}$, $\mathbf{y}$ należą do przestrzeni R^2 . Przestrzeń transformat jest także dwuwymiarowa, wektory $\mathbf{T}_h\mathbf{x}$, $\mathbf{T}_h\mathbf{y}$ należą do tej przestrzeni. Wektory $\mathbf{PT}_h(\mathbf{x} - \bar{\mathbf{w}}_h)$, $\mathbf{PT}_h(\mathbf{y} - \bar{\mathbf{w}}_h)$ są wynikiem rzutowania $\mathbf{P}$ wektorów $\mathbf{T}_h(\mathbf{x} - \bar{\mathbf{w}}_h)$, $\mathbf{T}_h(\mathbf{y} - \bar{\mathbf{w}}_h)$ na jednowymiarową podprzestrzeń wzorca (na rysunku 2(b) jest to prosta).



Rys. 2. (a) Rzutowanie w przestrzeni obserwacji.
(b) Rzutowanie w przestrzeni transformat Karhunen-Loève'a

Obliczenia wygodniej jest przeprowadzać w zredukowanej przestrzeni R^{M_h} . Wektor $\mathbf{z}$ tej przestrzeni uzyskuje się dla zadanego $\mathbf{x} \in R^L$ przez pozostawienie pierwszych M_h współrzędnych wektora $\mathbf{PT}_h\mathbf{x} \in R^L$, co zapiszemy:

$\mathbf{z} = r(\mathbf{PT}_h \mathbf{x})$. Macierz kowariancji $\mathbf{Q}_h$ wektorów uzyskanych jako wynik omówionych przekształceń wektorów $\mathbf{x} = \mathbf{w}_k - \bar{\mathbf{w}}_h$ dla $k \in W_h$ jest następująca:

$$\mathbf{Q}_h = \text{diag}(\lambda_1(\mathbf{R}_h), \lambda_2(\mathbf{R}_h), \dots, \lambda_{M_h}(\mathbf{R}_h)) \quad (18)$$

Macierz ta, zgodnie z uczynionym wcześniej założeniem, jest dodatnio określona. Odległość Mahalanobisa wektorów $\mathbf{x}$, $\mathbf{y}$ przestrzeni R^{M_h} zdefiniujemy następująco:

$$d_h(\mathbf{x}, \mathbf{y}) = \sqrt{(\mathbf{x} - \mathbf{y})^T \mathbf{Q}_h^{-1} (\mathbf{x} - \mathbf{y})} \quad (19)$$

6. Metryki dopasowane do wzorca W_h

W niniejszym punkcie proponujemy rozwiązanie problemu oceny stopnia podobieństwa wektora $\mathbf{x}$ przestrzeni obserwacji R^L do wzorca W_h . Rozwiązanie to bazuje na wyznaczonej wcześniej podprzestrzeni wzorca. Podstawą naszych propozycji są dwie metody wyznaczania odległości dopasowanej do wzorca W_h .

6.1. Metoda rzutowania na podprzestrzeń wzorca

Dopasowaną do wzorca W_h odległość między pomiędzy wektorami $\mathbf{x}$, $\mathbf{y}$ przestrzeni cech R^L proponujemy wyznaczać jako odległość między rzutami tych wektorów na podprzestrzeń wzorca. Biorąc pod uwagę fakt, że transformacja $\mathbf{T}_h$ Karhunen–Loève’a jest przekształceniem ortogonalnym, odpowiednie obliczenia można wykonać w przestrzeni transformat. Po wyznaczeniu rzutów $\mathbf{x}_h = r(\mathbf{PT}_h \mathbf{x})$, $\mathbf{y}_h = r(\mathbf{PT}_h \mathbf{y})$ w podprzestrzeni R^{M_h} poszukiwaną odległość oblicza się ze wzoru:

$$d_h(\mathbf{x}, \mathbf{y}) = d_h(\mathbf{x}_h, \mathbf{y}_h) = \sqrt{(\mathbf{x}_h - \mathbf{y}_h)^T \mathbf{Q}_h^{-1} (\mathbf{x}_h - \mathbf{y}_h)}, \quad \mathbf{x}, \mathbf{y} \in R^L \quad (20)$$

Metryka ta może być wykorzystywana do wyznaczania odległości między obserwacją $\mathbf{x}$ a klastrem $C(W_h)$. Przykładowo, odległość między obserwacją $\mathbf{x}$ a klastrem $C(W_h)$ dla metody centroidalnej jest określona wzorem:

$$D_h^{(1)}(\mathbf{x}, C(W_h)) = d_h(\mathbf{x}_h, \bar{\mathbf{w}}_h) = \sqrt{(\mathbf{x}_h - \bar{\mathbf{w}}_h)^T \mathbf{Q}_h^{-1} (\mathbf{x}_h - \bar{\mathbf{w}}_h)} \quad (21)$$

gdzie: $\mathbf{x}_h = r(\mathbf{PT}_h \mathbf{x})$, $\bar{\mathbf{w}}_h = r(\mathbf{PT}_h \bar{\mathbf{w}}_h)$, a wektor $\bar{\mathbf{w}}_h$ jest zdefiniowany przez (9).

6.2. Metoda obliczania odległości wektora obserwacji od podprzestrzeni wzorca

Metoda ta polega na bezpośrednim obliczaniu odległości między obserwacją $\mathbf{x}$ a klastrem $C(W_h)$. Odległość ta jest wyznaczana w przestrzeni cech R^L jako odległość między obserwacją $\mathbf{x}$ a jej rzutem $\mathbf{x}_h$ na podprzestrzeń wzorca $C(W_h)$:

$$D_h^{(2)}(\mathbf{x}, C(W_h)) = d(\mathbf{x}, \mathbf{x}_h) = \sqrt{(\mathbf{x} - \mathbf{x}_h)^T \mathbf{R}^{-1} (\mathbf{x} - \mathbf{x}_h)} \quad (22)$$

gdzie:

$$\mathbf{x}_h = \mathbf{T}_h^{-1} \mathbf{Pz} + \bar{\mathbf{w}}_h \quad (23)$$

$$\mathbf{z} = \begin{bmatrix} z_1 \\ \dots \\ z_{M_h} \\ z_{M_h+1} \\ \dots \\ z_L \end{bmatrix} = \mathbf{T}_h (\mathbf{x} - \bar{\mathbf{w}}_h), \quad \mathbf{Pz} = \begin{bmatrix} z_1 \\ \dots \\ z_{M_h} \\ 0 \\ \dots \\ 0 \end{bmatrix} \quad (24)$$

6.3. Porównywanie odległości

Porównywanie odległości dopasowanych do różnych wzorców wymaga normalizacji. Jej celem jest spełnienie dla każdego wzorca następującego warunku: wartość oczekiwana znormalizowanej odległości $\bar{D}_h^{(i)}(\mathbf{x}, C(W_h))$ ma wartość 1 [4]. Obliczenie wartości znormalizowanej określa następujący wzór:

$$\bar{D}_h^{(i)}(\mathbf{x}, C(W_h)) = \frac{1}{\sqrt{N_h}} D_h^{(i)}(\mathbf{x}, C(W_h)), \quad i = 1, 2 \quad (25)$$

6.4. Skalaryzacja wektora odległości

Wykorzystywanie obydwu metod oceny odległości analizowanego punktu $\mathbf{x}$ przestrzeni obserwacji R^L od wzorca W_h daje możliwość wektorowej oceny tej odległości; jej wynikiem są odległości $\bar{D}_h^{(1)}(\mathbf{x}, C(W_h))$ oraz $\bar{D}_h^{(2)}(\mathbf{x}, C(W_h))$.

Ocena druga przyjmuje wartość zerową wtedy, gdy podprzestrzeń wzorca nie jest właściwa, tzn. gdy $R^{M_h} = R^L$. Łatwo stwierdzamy, że jeśli $M_h = L$, to $\mathbf{Pz} = \mathbf{z}$. W konsekwencji $\mathbf{x} - \mathbf{x}_h = \mathbf{x} - \mathbf{T}_h^{-1}\mathbf{Pz} - \bar{\mathbf{w}}_h = \mathbf{x} - \mathbf{T}_h^{-1}\mathbf{T}_h(\mathbf{x} - \bar{\mathbf{w}}_h) - \bar{\mathbf{w}}_h = \mathbf{0}$. Wtedy $D_h^{(2)}(\mathbf{x}, C(W_h)) = 0$ dla wszystkich $\mathbf{x} \in R^L$.

Liniowe uporządkowanie według wektorowych obliczeń odległości można uzyskać, stosując dowolną metodę skalaryzacji wektora odległości. Przykładowo, zastosowanie wzoru:

$$\bar{D}_h^s(\mathbf{x}, C(W_h)) = \sqrt{[\bar{D}_h^{(1)}(\mathbf{x}, C(W_h))]^2 + [\bar{D}_h^{(2)}(\mathbf{x}, C(W_h))]^2} \quad (26)$$

umożliwia „gładkie” przejście od oceny z dominacją wartości $\bar{D}_h^{(2)}(\mathbf{x}, C(W_h))$ do oceny wyłącznie na podstawie wartości $\bar{D}_h^{(1)}(\mathbf{x}, C(W_h))$. Sytuacja taka może mieć miejsce w przypadku wzbogacania wzorca W_h przez wskazywanie nowych jego przykładów.

7. Eksperyment

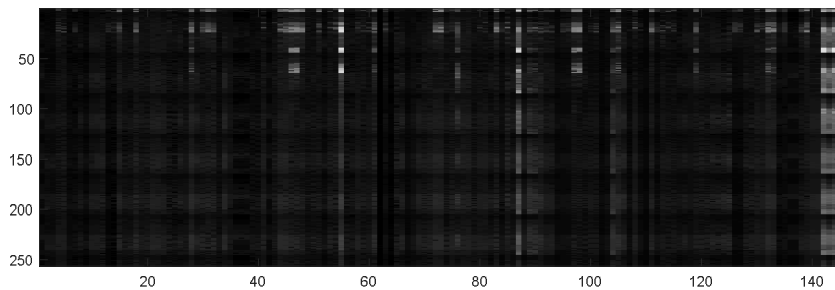
7.1. Przedmiot i cel badań

W celu zilustrowania proponowanych metod dokonano analizy przykładowego zbioru wyników pomiaru¹. Pojedynczy wynik pomiaru stanowi wektor, którego współrzędne były wyznaczone jako wyniki benchmarków. Każdy pomiar wykonywany był na innym zestawie komputerowym, współrzędne o jednakowym indeksie opisują wynik tego samego benchmarku. Zestawy miały różną konfigurację sprzętową i programową, tzn. różniły się albo procesorami, albo płytami głównymi, albo systemami operacyjnymi, albo zainstalowanym, aktywnym oprogramowaniem i otoczeniem sieciowym. Wykorzystywany zbiór danych zawierał wyniki 256 benchmarków wyznaczone dla 145 zestawów. Wizualizację tych wyników pomiaru w postaci obrazu przedstawiono na rysunku 3.

Istotną cechą zastosowanej metody analizy jest brak wymagania znajomości charakterystyk analizowanego zbioru danych – zadanego po prostu w postaci macierzy. Jako wzorzec zachowania normalnego wskazano 19 zestawów bazujących na procesorze typu A. Jako wzorzec alternatywny zostało

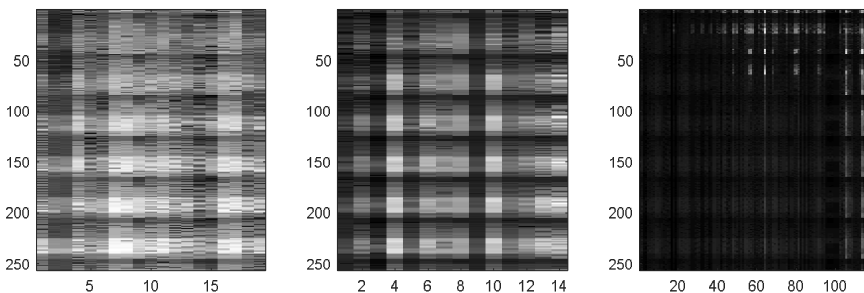
¹ Do obliczeń zostały wykorzystane wyniki pomiarów udostępnione mi przez ich autora, Artura Miktusa (artur.miktus@wat.edu.pl).

wskazanych 14 zestawów bazujących na procesorach typu B. Przy tak wybranych wzorcach analiza polegała na określeniu, które z badanych zestawów zachowują się jak zestawy wzorcowe wyposażone w procesory typu A, a które jak zestawy wzorcowe wyposażone w procesory typu B.



Rys. 3. Wizualizacja źródłowych wyników pomiaru w postaci obrazu. Wyniki pomiaru są liczbami dodatnimi. Stopień szarości odpowiada wartości liczbowej elementu macierzy. Liczba wierszy jest równa liczbie współrzędnych wektora pomiaru 256, liczba kolumn jest równa liczbie zestawów 145

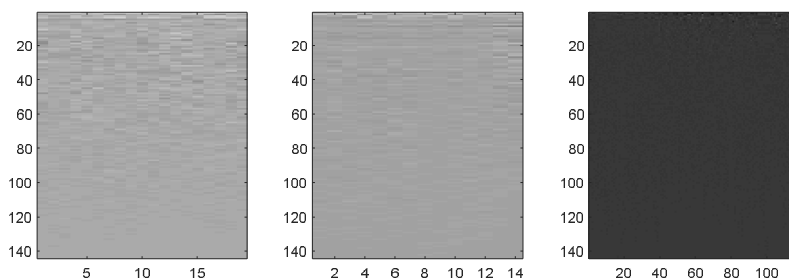
Do testowania wybrano arbitralnie 117 zestawów, dla których wskazany przykładami rodzaj anomalii nie był oczekiwany. W zbiorze wybranych do testowania zestawów umieszczono też zestawy wskazane jako wzorce zachowania normalnego. Wizualizacja macierzy pomiarów dla wskazanych zestawów wzorcowych oraz zestawów testowanych została przedstawiona na rysunku 4.



Rys. 4. Wizualizacja macierzy pomiarów: na rysunku z lewej dla wzorca W_1 , na rysunku w środku dla wzorca W_2 oraz na rysunku z prawej dla zestawów testowanych. Liczba przykładów wzorca W_1 jest równa $N_1 = 19$, liczba przykładów wzorca W_2 jest równa $N_2 = 14$, liczba zestawów testowanych $N_a = 117$

7.2. Przetwarzanie wstępne

Celem wstępnego przetwarzania analizowanych danych była redukcja liczby pomiarów, tak aby odpowiadająca im macierz kowariancji była dodatnio określona. Rezultat ten został osiągnięty poprzez wyznaczenie składowych głównych o niezerowej wariancji. Liczba takich składowych wyniosła 144. Uzyskana macierz składowych głównych była podstawą dalszej analizy i nazywana jest dalej macierzą wyników obserwacji. Odpowiednie wartości parametrów w przedstawianym przykładzie są więc następujące: liczba wszystkich przykładów zachowań $N=145$, liczba wskazanych przykładów dla wzorca pierwszego $N_1=19$, liczba wskazanych przykładów dla wzorca drugiego $N_2=14$, liczba współrzędnych wektora obserwacji $L=144$. Przedmiotem analizy jest więc 145 wektorów obserwacji, a każdy wektor obserwacji ma 144 współrzędne. Wizualizację obliczonych wyników obserwacji dla wskazanych przykładów wzorców oraz testowanych zestawów przedstawiono na rysunku 5.

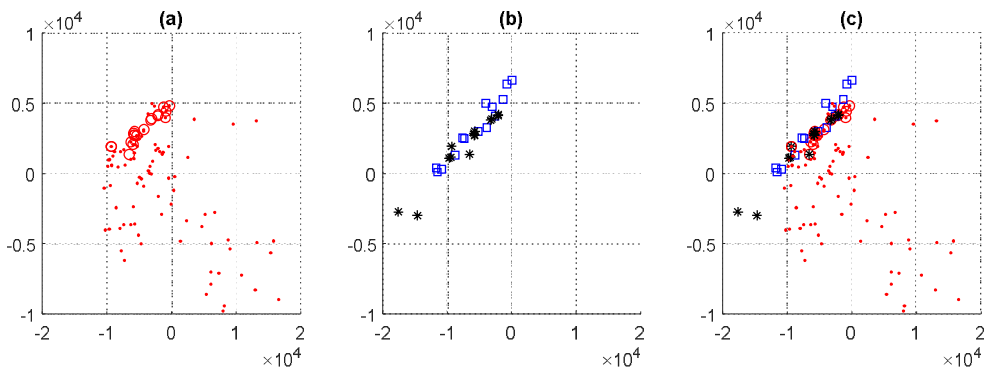


Rys. 5. Wizualizacja macierzy obserwacji: na rysunku z lewej dla wzorca W_1 , na rysunku w środku dla przykładów wzorca W_2 oraz na rysunku z prawej dla testowanych zestawów. Liczba współrzędnych wektora składowych głównych jest równa $L=144$, liczba przykładów wzorca W_1 jest równa $N_1=19$, liczba przykładów wzorca W_2 jest równa $N_2=14$, liczba testowanych zestawów jest równa $N_a=117$

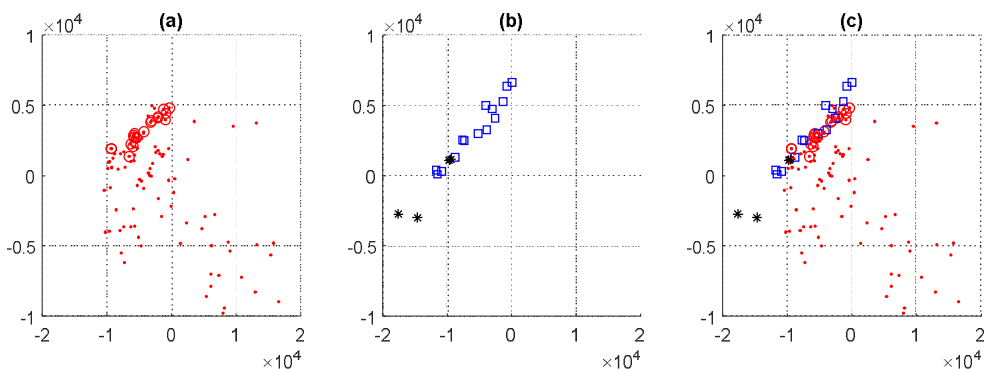
7.3. Wyniki analizy

W przedstawianym przykładzie podprzestrzeń wzorca W_1 ma wymiar $M_1=18$, a podprzestrzeń wzorca W_2 ma wymiar $M_2=13$. Porównanie odległości badanego wektora obserwacji $\mathbf{x} \in R^{144}$ od klastra $C(W_1)$ wzorca W_1 i odległości badanego wektora obserwacji $\mathbf{x} \in R^{144}$ od klastra $C(W_2)$ wzorca

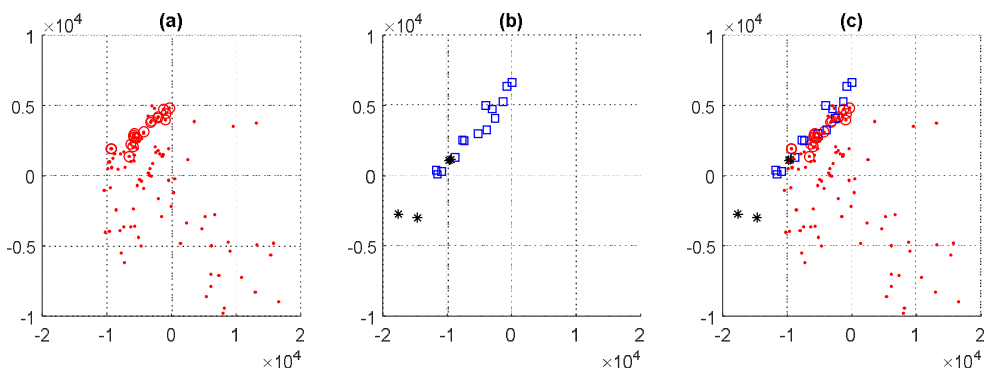
W_2 zostało wykonane dla odległości $\overline{D}_h^{(1)}(\mathbf{x}, C(W_h))$, $\overline{D}_h^{(2)}(\mathbf{x}, C(W_h))$ oraz $\overline{D}_h^s(\mathbf{x}, C(W_h))$. Wyniki porównań zostały zobrazowane na rysunkach 6-11.



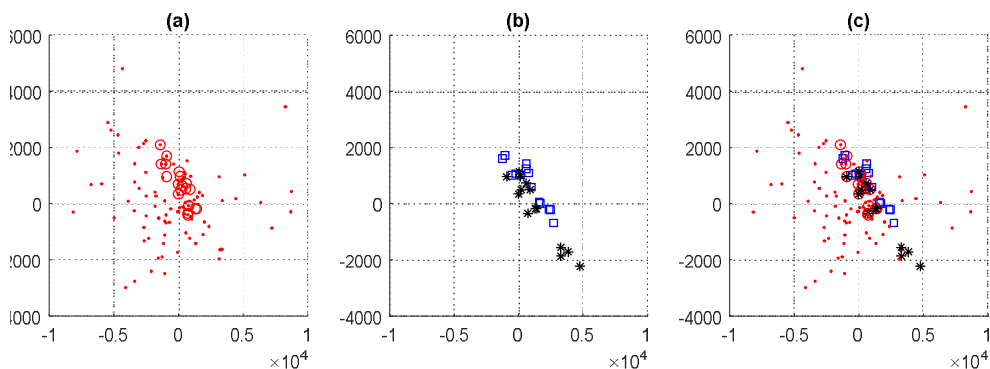
Rys. 6. Wizualizacja wyników analizy w podprzestrzeni pierwszych dwóch współrzędnych wektorów obserwacji na podstawie porównania odległości $\overline{D}_h^{(1)}(\mathbf{x}, C(W_h))$. (a) Wzorce normalności W_1 są oznaczone kółkami, punktami oznaczono wyniki obserwacji ocenione jako normalne. (b) Wzorce anomalii W_2 są oznaczone kwadratami, gwiazdkami oznaczono wyniki obserwacji ocenione jako anomalne. (c) Nałożenie wykresów (a) oraz (b)



Rys. 7. Wizualizacja wyników analizy w podprzestrzeni pierwszych dwóch współrzędnych wektorów obserwacji na podstawie porównania odległości $\overline{D}_h^{(2)}(\mathbf{x}, C(W_h))$. (a) Wzorce normalności W_1 są oznaczone kółkami, punktami oznaczono wyniki obserwacji ocenione jako normalne. (b) Wzorce anomalii W_2 są oznaczone kwadratami, gwiazdkami oznaczono wyniki obserwacji ocenione jako anomalne. (c) Nałożenie wykresów (a) oraz (b)



Rys. 8. Wizualizacja wyników analizy w podprzestrzeni pierwszych dwóch współrzędnych wektorów obserwacji na podstawie porównania odległości $\bar{D}_h^s(x, C(W_h))$. (a) Wzorce normalności W_1 są oznaczone kółkami, punktami oznaczono wyniki obserwacji ocenione jako normalne. (b) Wzorce anomalii W_2 są oznaczone kwadratami, gwiazdkami oznaczono wyniki obserwacji ocenione jako anomalne. (c) Nałożenie wykresów (a) oraz (b)

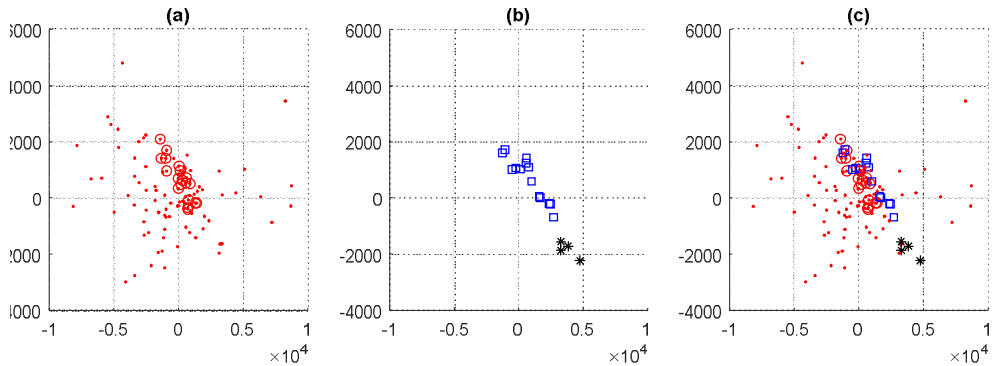


Rys. 9. Wizualizacja wyników analizy w podprzestrzeni współrzędnych o indeksach 3 i 4 wektorów obserwacji na podstawie porównania odległości $\bar{D}_h^{(1)}(x, C(W_h))$. (a) Wzorce normalności W_1 są oznaczone kółkami, punktami oznaczono wyniki obserwacji ocenione jako normalne. (b) Wzorce anomalii W_2 są oznaczone kwadratami, gwiazdkami oznaczono wyniki obserwacji ocenione jako anomalne. (c) Nałożenie wykresów (a) oraz (b)

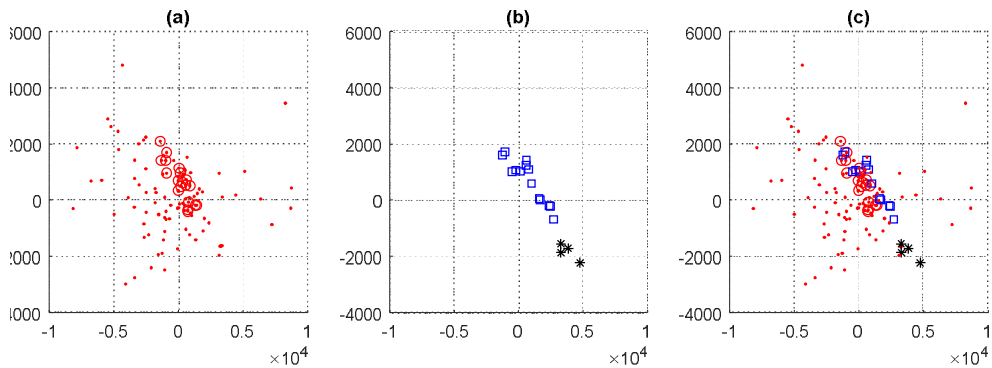
Ostateczny wynik obliczeń stanowi wskazanie, które z badanych zestawów były bliższe wzorcowi normalności, a które – wzorcowi alternatywnemu, interpretowanemu jako wzorec anomalii. Na podstawie analizy odległości w podprzestrzeniach wzorców wykryto 14 przypadków anomalii (por. rys. 6 i rys. 9). Na podstawie analizy odległości obserwacji od podprzestrzeni wzorca

wykryto 4 przypadki (rys. 7 i 10), które pokrywają się z wykryciami na podstawie odległości sumarycznej (rys. 8 i 11).

Uzyskane wskazania anomalii mogą stanowić podstawę dodatkowej, szczegółowej analizy, mającej na celu przedstawienie przyczyn anomalii. W przeprowadzonym eksperymencie taka dodatkowa analiza była możliwa, ponieważ wszystkie zestawy i pomiary były dokładnie opisane.



Rys. 10. Wizualizacja wyników analizy w podprzestrzeni współrzędnych o indeksach 3 i 4 wektorów obserwacji na podstawie porównania odległości $\overline{D}_h^{(2)}(x, C(W_h))$. (a) Wzorce normalności W_1 są oznaczone kółkami, punktami oznaczono wyniki obserwacji ocenione jako normalne. (b) Wzorce anomalii W_2 są oznaczone kwadratami, gwiazdkami oznaczono wyniki obserwacji ocenione jako anomalne. (c) Nałożenie wykresów (a) oraz (b)



Rys. 11. Wizualizacja wyników analizy w podprzestrzeni współrzędnych o indeksach 3 i 4 wektorów obserwacji na podstawie porównania odległości $\overline{D}_h^s(x, C(W_h))$. (a) Wzorce normalności W_1 są oznaczone kółkami, punktami oznaczono wyniki obserwacji ocenione jako normalne. (b) Wzorce anomalii W_2 są oznaczone kwadratami, gwiazdkami oznaczono wyniki obserwacji ocenione jako anomalne. (c) Nałożenie wykresów (a) oraz (b)

Szczegółowa analiza opisu zestawów i warunków pomiarów (ustalająca przyczyny uzyskania obserwacji odbiegających od wzorcowych) może być podstawą orzeczenia o wystąpieniu ewentualnego błędu detekcji: fałszywego alarmu lub fałszywego spokoju, a także przedstawienia wniosków o zasadach optymalizacji dwukryterialnej (np. przez określenie ważonej odległości sumarycznej).

8. Wnioski końcowe

- 1) Proponowane metody detekcji anomalii mają charakter uniwersalny i mogą być wykorzystywane wszędzie tam, gdzie wyniki zachowania systemu można traktować jako pomiary zapisywane w postaci wektorów liczb rzeczywistych. Nie jest potrzebne opracowanie żadnych modeli powstawania wykorzystywanych wyników pomiarów.
- 2) Definiowanie normalności sprowadza się do wskazania odpowiednich przykładów. Wskazanie przykładów anomalii pozwala uprościć wnioskowanie dzięki możliwości porównywania dwóch odległości obserwowanego zachowania systemu: od wskazanych przykładów zachowania normalnego i analogicznej odległości od wskazanych przykładów zachowania anomального. Wynikiem wnioskowania jest wskazanie systemów, dla których wyniki obserwacji ich zachowań odbiegają od wskazanych przykładów zachowań wzorcowych.
- 3) Przedstawione dwie zasadnicze metody obliczania odległości mogą być wykorzystywane dowolnie, w zależności od badanego systemu. Zastosowanie łączne prowadzi do optymalizacji wektorowej. W przypadku wskazania wystarczająco dużo przykładów wzorca (w stosunku do wymiaru wektora obserwacji) odległość analizowanej obserwacji od podprzestrzeni wzorca staje się zerowa i w ten sposób uzyskuje się płynne przejście do oceny odległości tylko na podstawie metryki w przestrzeni wzorca.
- 4) Przedstawione metody obliczeniowe pozwalają na analizę w sytuacji, gdy wymiar wektora pomiaru jest większy od liczby wskazanych przykładów wzorca. Z reguły dotyczy to wzorców anomalii. Sytuacja taka występuje praktycznie wtedy, gdy wyniki obserwacji systemu nie są selekcjonowane pod kątem ich użyteczności w zadaniach detekcji anomalii. Dotyczy to zwłaszcza zadań wykrywania anomalii na podstawie danych generowanych automatycznie, zwykle przeznaczonych do innych celów.
- 5) Tworzenie wzorców przez wskazywanie przykładów można traktować jako sposób łączenia (syntezy, fuzji) informacji pochodzących z różnych źródeł.

Łączenie to ma charakter sekwencyjny w tym sensie, że wskazywane przykłady są uzyskiwane na podstawie obserwacji z innych źródeł.

Literatura

- [1] CAMPOS G.O., ZIMEK A., SANDER J., CAMPELLO R.J.G.B., MICENKOVÁ B., SCHUBERT E., ASSENT I., HOULE M.E., *On the evaluation of unsupervised outlier detection: measures, datasets, and an empirical study*. Data Mining and Knowledge Discovery 30(4), 2016, pp. 891-927.
- [2] CHANDOLA V., BANERJEE A., KUMAR V., *Anomaly detection: A survey*. ACM Computing Surveys, Vol. 41, No. 3, Article 15, 2009.
- [3] HODGE V.J., AUSTIN J., *A survey of outlier detection methodologies*. Artificial Intelligence Review, 22 (2), 2004, pp. 85-126.
- [4] KWIATKOWSKI W., *Metody automatycznego rozpoznawania wzorców*, BEL Studio, Warszawa, 2010.
- [5] SODEMANN A.A., ROSS M.P., BORGHETTI B.I., *A review of anomaly detection in automated surveillance*. IEEE Transactions on Systems Man and Cybernetics Part C (Applications and Reviews), Vol. 42, No. 6, 2012, 1257-1272.
- [6] PIMENTEL M.A., CLIFTON D.A., CLIFTON L., TARASSENKO L., *A review of novelty detection*. Signal Processing, Vol. 99, 2014, 215-249.
- [7] TUKEY J.W., *Exploratory data analysis*. Addison-Wesley, 1977.
- [8] YAO-GUANG WEI, DE-LING ZHENG, YING WANG, *Research of a negative selection algorithm and its application in anomaly detection*. Proceedings of 2004 International Conference on Machine Learning and Cybernetics (IEEE Cat. No.04EX826), Vol. 5, 2004, 2910-2913.
- [9] ZIMEK A., SCHUBERT E., KRIEGEL H., *A survey on unsupervised outlier detection in high-dimensional numerical data*. Statistical Analysis and Data Mining, Vol. 5, Issue 5, 2012, 363-387.

Anomaly detection based on given examples

ABSTRACT: The paper considers the issue of anomalies detection based on registered observations of a system behavior. The problem is formulated as recognition of normal and anomalous behavior patterns. Both types of patterns are identified by indication of appropriate examples. A peculiarity of this task is that usually the number of examples is far lower than the dimension of vectors describing the observations. Two methods to solve this task have been presented in the paper, based on projecting the observations on the subspace of examples. The first method is based on a distance of the observation vector from the subspace of examples. The second method is based on transferring the pattern recognition problem to the subspace of examples.

KEYWORDS: anomaly detection, novelty detection, outlier detection, pattern recognition, exploratory data analysis, Mahalanobis distance

Praca wpłynęła do redakcji: 18.05.2018 r.

Koncepcja zwiększenia bezpieczeństwa transakcyjnych aplikacji internetowych

Łukasz STRZELECKI, Andrzej LUDEW

Instytut Teleinformatyki i Automatyki, Wydział Cybernetyki WAT,
ul. Gen. W. Urbanowicza 2, 00-908 Warszawa
lukasz.strzelecki@wat.edu.pl, andrzej.ludew@wat.edu.pl

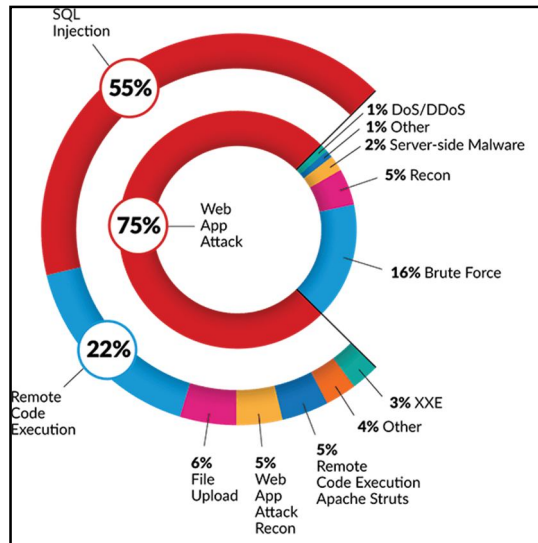
STRESZCZENIE: W artykule przedstawiono nowatorską koncepcję podniesienia bezpieczeństwa istotnych aplikacji internetowych wykorzystywanych np. podczas realizacji operacji finansowych, poprzez niestandardowe wykorzystanie mechanizmów języka HTML5 oraz tzw. *przetwarzania po stronie serwera*. Omówiono zalety proponowanego podejścia oraz wskazano jego potencjalne wady.

SŁOWA KLUCZOWE: bezpieczeństwo, witryny WWW, aplikacje internetowe, HTML5, Canvas

1. Wprowadzenie

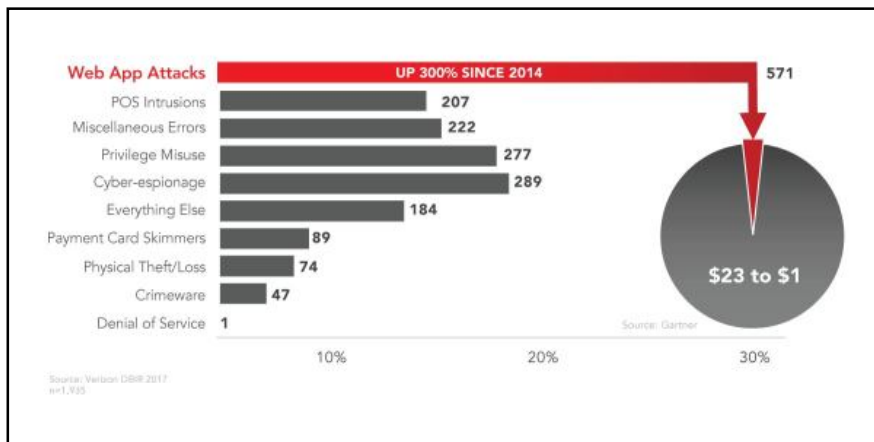
Kwestie związane z bezpieczeństwem teleinformatycznym oraz stale wzrastającą liczbą odnotowywanych incydentów teleinformatycznych są często poruszane w literaturze przedmiotu. Warto jednakże zwrócić uwagę, że zdecydowana większość spośród nich związana jest z tzw. *aplikacjami internetowymi*, rozumianymi jako interaktywne oprogramowanie działające po stronie serwera, komunikujące się z użytkownikami przy wykorzystaniu przeglądarek internetowych. Według danych firmy *Alert Logic*, oferującej rozwiązanie typu SOC, ataki na aplikacje internetowe stanowiły 75% wszystkich incydentów odnotowanych przez oferowane przez nią rozwiązania¹.

¹ Źródło: <https://www.alertlogic.com/solutions/activewatch-managed-detection-and-response/>



Rys. 1. Typy ataków wykryte przez rozwiązania firmy Alert Logic
(źródło: <https://alertlogic.com>)

Warto dodać, iż wspomniana firma powołuje się również na raport *2017 Data Breach Investigations Report*² opracowany przez firmę Verizon. Zgodnie z informacjami zawartymi we wspomnianym dokumencie, udział ataków na aplikacje internetowe stale wzrasta, a od 2014 roku zwiększył się ponad 3-krotnie.



Rys. 2. Zestawienie typów ataków odnotowanych przez firmę Verizon
(źródło: <https://alertlogic.com>)

² Źródła: <https://www.ictsecuritymagazine.com/wp-content/uploads/2017-Data-Breach-Investigations-Report.pdf>;
<https://www.alertlogic.com/solutions/why-alert-logic/full-stack-security/>

Powyższe informacje potwierdzają również raporty *Krajobraz bezpieczeństwa Polskiego Internetu* opublikowany w 2016 roku przez CERT Polska³ oraz *Raport CERT Orange Polska za rok 2017*⁴, w których opisane są liczne incydenty związane z atakami na popularne aplikacje internetowe.

Uwzględniając powyższe, w niniejszym artykule zaproponowano nowatorską metodę budowy i zabezpieczania szczególnie wrażliwych aplikacji internetowych, za które należy uznać m.in. serwisy transakcyjne banków (tzw. *bankowość internetowa*) dostępne dla standardowych użytkowników. Koncentracja na tym konkretnym typie aplikacji internetowych wynika bezpośrednio z ich specyfiki (konieczność logowania się użytkowników) oraz istotności, która przekłada się bezpośrednio na środki, które poszczególne instytucje są w stanie poświęcić na zwiększenie ich bezpieczeństwa (np. poprzez ich przebudowę do proponowanego modelu działania). W dalszej części artykułu wspomniany typ aplikacji internetowych będzie nazywany *transakcyjnymi aplikacjami internetowymi*, natomiast zawartość renderowana i wyświetlana bezpośrednio w przeglądarce użytkownika będzie nazywana *witryną internetową* (niezależnie od tego, w jaki sposób jej kod został wytworzony⁵).

2. Ogólna budowa tradycyjnych aplikacji internetowych

W pewnym uproszczeniu, abstrahując od technologii wykorzystywanych zarówno po stronie serwera, jak i klienta można przyjąć, że:

- a) przeglądarka internetowa wysyła do serwera WWW asynchroniczne żądanie⁶,
- b) serwer odpowiadając na otrzymane żądanie, dostarcza przeglądarce internetowej użytkownika treść, która jest renderowana i wyświetlana

³ Raport CERT Polska jest dostępny pod adresem URL:

https://www.cert.pl/PDF/Raport_CP_2016.pdf

⁴ Raport CERT Orange Polska jest dostępny pod adresem URL:

<https://cert.orange.pl/opracowania>

⁵ To znaczy niezależnie od tego, czy bazuje ona na statycznych plikach HTML zamieszczonych na serwerze, czy jej kod HTML jest dynamicznie generowany przez (transakcyjną) aplikację internetową.

⁶ Protokół HTTP nie przewiduje utrzymywania sesji użytkownika ani rozróżniania sekwencji żądań i stąd wszystkie żądania wykonywane z jego wykorzystaniem mają charakter asynchroniczny. Warto natomiast nadmienić, że istnieją rozwiązania, które zapewniają utrzymywanie stanu sesji/widoku użytkownika, np. *ViewState*.

w przeglądarce klienta lub w inny sposób wpływa na działanie witryny wyświetlanej w oknie przeglądarki internetowej użytkownika⁷.

Wysyłanie żądania od klienta (przeglądarki internetowej) do serwera zazwyczaj spowodowane jest jednym z trzech podstawowych zdarzeń:

- a) użytkownik odwołał się do adresu URL zamieszczonego na wyświetlanej mu witrynie;
- b) skrypt języka *JavaScript* funkcjonujący w ramach witryny wyświetlanej użytkownikowi automatycznie wykonał żądanie (np. w celu zaktualizowania danych wyświetlanych użytkownikowi);
- c) użytkownik wypełnił i zatwierdził formularz HTML znajdujący się na wyświetlanej mu witrynie.

W każdym z powyższych przypadków żądanie wysyłane jest do serwera (przez przeglądarkę użytkownika) z wykorzystaniem odpowiedniej metody HTTP (m.in. GET, POST itp.). Ponadto, jeśli do przesłanego żądania dołączone były parametry (których oczekiwała aplikacja działająca po stronie serwera), to są one dodatkowo przetwarzane. Warto zwrócić uwagę, że to właśnie na tym etapie popełnianych jest przez programistów aplikacji internetowych najwięcej błędów⁸ i w związku z tym należy się spodziewać występowania tutaj największej liczby podatności na ataki. Ze względu na powyższe duża część testów penetracyjnych aplikacji internetowych koncentruje się na weryfikacji poprawności przetwarzania danych przesyłanych do serwera (i osadzonej na nim aplikacji internetowej). Należy też zauważyć, że sposób działania zdecydowanej większości tzw. *automatycznych skanerów bezpieczeństwa* przeznaczonych dla aplikacji internetowych sprowadza się głównie do symulowania działania przeglądarki internetowej użytkownika i przesyłania (podszywając się pod nią) do serwera różnego typu ciągów formatujących, które mają na celu zweryfikowanie, czy aplikacja działająca po stronie serwera poprawnie je obsłużyła (albo zwróciła komunikat pozwalający wnioskować o występowaniu luki bezpieczeństwa).

Analizując sposób działania tradycyjnych aplikacji internetowych, warto także zwrócić uwagę, że zazwyczaj po stronie serwera działa jeden program omawianego typu i samodzielnie obsługuje wszystkie żądania użytkowników do niego kierowane⁹. W celu rozróżnienia żądań kierowanych od poszczególnych użytkowników wykorzystywany jest mechanizm tzw. *ciasteczek*, czyli

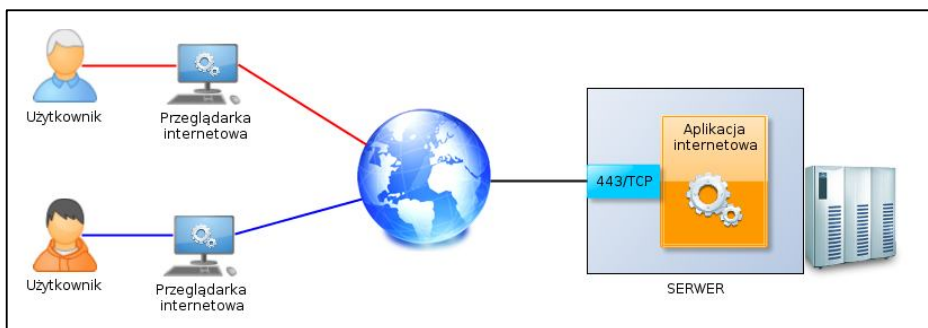
⁷ Opisany sposób działania znajduje zastosowanie zarówno w przypadku statycznych stron WWW, jak i w przypadku aplikacji internetowych korzystających np. z technologii *Asynchronous JavaScript and XML* (AJAX).

⁸ Do najczęstszych błędów należy brak odpowiedniego, bezpiecznego przetwarzania danych pochodzących od użytkownika, tj. przyjmowane jest często założenie, że dane przesłane do serwera są poprawne.

⁹ W przypadku aplikacji uruchamianych oddzielnie dla każdego żądania należy zwrócić uwagę, że w rzeczywistości sprowadza się to do działania tego samego kodu w tym samym środowisku, a zatem można tę sytuację interpretować jako działanie pojedynczej aplikacji.

niewielkich, unikatowych ciągów znaków zapisywanych w pamięci przeglądarki internetowej użytkownika. Wspomniany ciąg jest dołączany do każdego żądania wysyłanego przez przeglądarkę do serwera i na tej podstawie konkretny użytkownik może zostać zidentyfikowany. Dla porządku należy zaznaczyć, że wspomniane *ciasteczka* również są częstym celem ataków, gdyż ich przechwycenie umożliwia podszycie się pod (zalogowanego do danej witryny) użytkownika.

W celach poglądowych ogólny schemat sposobu działania tradycyjnych aplikacji internetowych został przedstawiony na rysunku 3.



Rys. 3. Schemat poglądowy sposobu działania tradycyjnych aplikacji internetowych (źródło: opracowanie własne)¹⁰

3. Koncepcja zwiększenia bezpieczeństwa transakcyjnych aplikacji internetowych

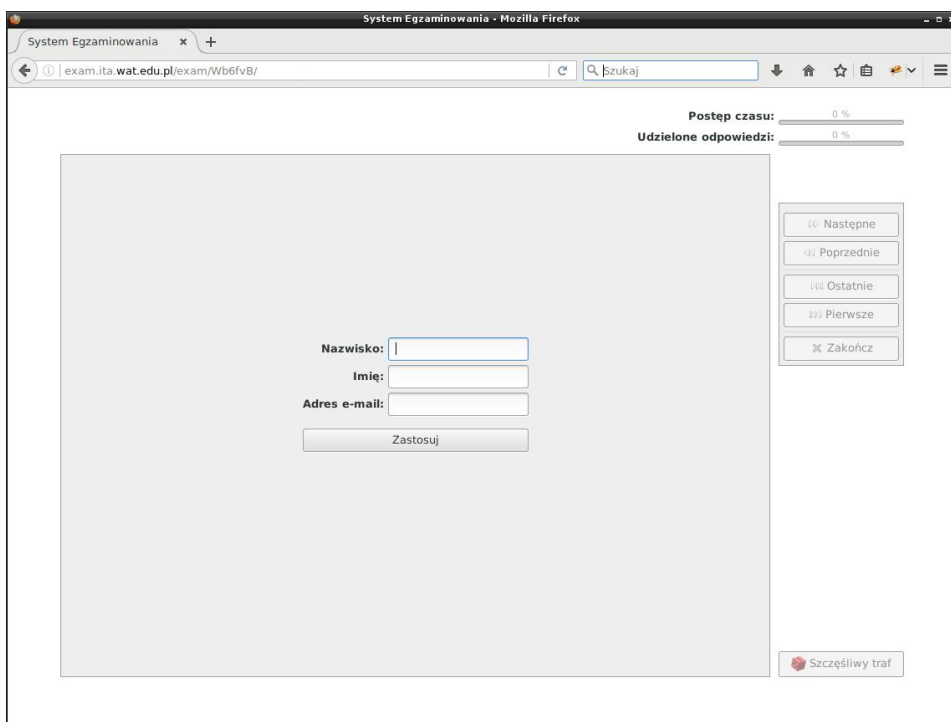
Wraz z wprowadzeniem w 2014 roku standardu HTML 5 zestaw elementów możliwych do wykorzystania w ramach witryn internetowych został rozbudowany m.in. o znacznik `<canvas>`. Jego podstawowym zastosowaniem – zgodnie z założeniami projektantów – było umożliwienie rysowania oraz tworzenia interaktywnych wykresów bezpośrednio na stronie internetowej. Jednakże praktyka pokazuje, że zakres jego zastosowań jest istotnie większy, gdyż umożliwia on również:

- wyświetlanie filmów/animacji,
- wyświetlanie obrazu interaktywnych gier,

co z kolei umożliwia wykorzystanie go także jako ekranu dla stacjonarnej

¹⁰ Na rysunku przedstawiono (w celach poglądowych) sytuację, w której serwer nasłuchuje na porcie 443/TCP z tego względu, iż port ten zazwyczaj jest wykorzystywany przy połączeniach szyfrowanych, a te z kolei są standardem w przypadku *transakcyjnych aplikacji internetowych*.

aplikacji (tj. takiej, która nie została przewidziana do pracy w formie aplikacji internetowej). Należy podkreślić, że w takim przypadku aplikacja pozostanie funkcjonalna, a użytkownik będzie mógł w niej wykonywać wszystkie operacje, jednakże z wykorzystaniem przeglądarki internetowej, w ramach której będzie wyświetlane jej okno. Taki sposób działania aplikacji graficznych dopuszcza między innymi biblioteka GTK3 – na rysunku 4 przedstawiony został przykład okna programu stacjonarnego wyświetlonego w ramach przeglądarki internetowej *Firefox*.



Rys. 4. Widok aplikacji stacjonarnej *System Egzaminowania* uruchomionej w ramach przeglądarki internetowej *Firefox* (źródło: opracowanie własne)

Należy w tym miejscu nadmienić, że zmiana sposobu wyświetlania głównego okna programu (tj. wyświetlanie go w ramach przeglądarki internetowej) nie powoduje jego automatycznej transformacji w typową aplikację internetową. Jest tak m.in. dlatego, że do jego interakcji z użytkownikiem (np. przy wprowadzaniu i zatwierdzaniu danych) nie są stosowane formularze HTML jak w tradycyjnych aplikacjach internetowych, lecz połączenia wykorzystujące mechanizm *WebSocket*¹¹, za pośrednictwem którego przesyłane są dane binarne

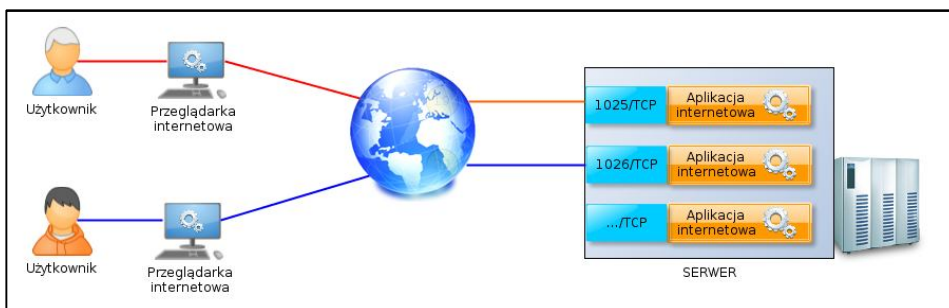
¹¹ Mechanizm *WebSocket* jest zdefiniowany w dokumencie normatywnym RFC 6455 i od roku 2011 jest obsługiwany przez wszystkie popularne przeglądarki sieci Internet.

informującą aplikację m.in. o tym, że w ramach jej okna użytkownik poruszył kursorem myszy, wcisnął przycisk, czy też wprowadził ciąg znaków. Warto podkreślić, że zmiana sposobu przesyłania danych do aplikacji stanowi główny mechanizm zwiększający jej bezpieczeństwo, co zostanie rozwinięte w kolejnym punkcie.

Analizując sposób działania aplikacji stacjonarnej, której okno główne wyświetlane jest w ramach przeglądarki internetowej, konieczne trzeba zwrócić uwagę, iż ze względu na specyfikę mechanizmu *WebSocket*:

- a) tylko jeden klient w danym momencie (poprzez przeglądarkę internetową) może korzystać z konkretnej instancji aplikacji stacjonarnej – każda próba dołączenia nowego klienta do już wykorzystanego gniazda *WebSocket* spowoduje rozłączenie klienta wcześniej połączonego;
- b) liczba aplikacji obsługujących klientów musi być nie mniejsza niż ich liczba,
- c) każde połączenie wymaga wykorzystania innego portu TCP po stronie serwera.

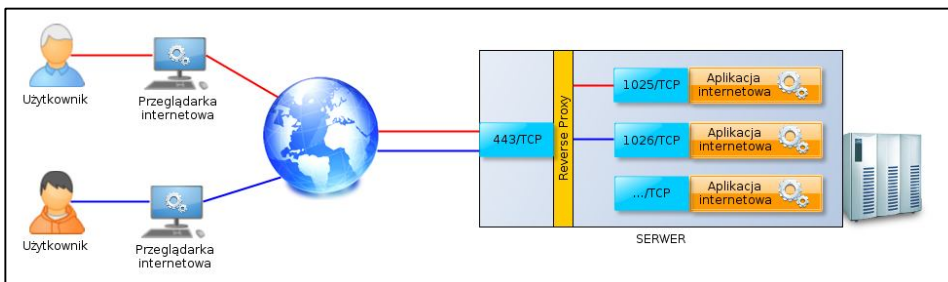
Powyższe ograniczenia powodują, że w ogólności schemat połączenia klientów (za pomocą przeglądarek sieci Internet) do aplikacji stacjonarnych udostępnianych z wykorzystaniem elementu *Canvas* języka HTML5 musi być analogiczny do przedstawionego na rysunku 5.



Rys. 5. Schemat poglądowy proponowanego sposobu działania transakcyjnych aplikacji internetowych (źródło: opracowanie własne)

Łatwo zauważyć, że powyżej przedstawiony schemat działania jest niedopuszczalny w zdecydowanej większości przypadków, gdyż zazwyczaj połączenia do aplikacji internetowych są dozwolone jedynie przy wykorzystaniu portów serwera 80/TCP i 443/TCP (w przypadku połączeń szyfrowanych). Ze względów bezpieczeństwa połączenia do innych portów są blokowane bezpośrednio na zaporach sieciowych (ochraniających serwery udostępniające aplikacje internetowe). Jednakże analiza technicznych możliwości ominięcia opisanego problemu wskazuje, że rozwiązaniem może być zastosowanie mechanizmu tzw. *Reverse Proxy*. Umożliwia on mianowicie udostępnianie wielu

aplikacji internetowych poprzez ten sam port, jednakże przy wykorzystaniu odmiennych adresów URL. Ogólny schemat podłączenia użytkowników do w ten sposób udostępnianych aplikacji stacjonarnych (przy wykorzystaniu mechanizmów *WebSocket* oraz *Reverse Proxy*) został przedstawiony na rysunku 6.

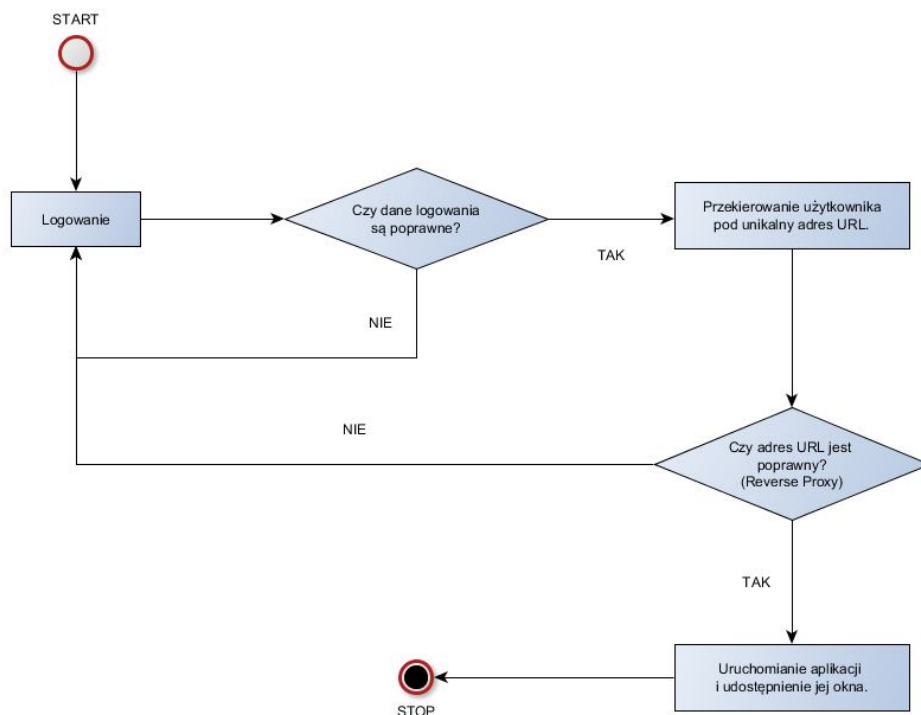


Rys. 6. Schemat poglądowy proponowanego sposobu działania transakcyjnych aplikacji internetowych (źródło: opracowanie własne)

Reasumując, jako praktycznie realizowalny sposób zwiększenia bezpieczeństwa *transakcyjnych aplikacji internetowych* można przyjąć ich przygotowanie w formie aplikacji stacjonarnych (z założenia odpornych na zdecydowaną większość ataków ukierunkowanych na tradycyjne aplikacje internetowe), których główne okna będą udostępniane użytkownikom przy wykorzystaniu mechanizmu *WebSocket* ukrytego za mechanizmem *Reverse Proxy*. W tej sytuacji standardowy sposób obsługi każdego użytkownika powinien sprowadzać się do:

- odwiedzenia przez niego strony głównej wybranej instytucji;
- skierowania go do aplikacji internetowej, która – po jego poprawnym uwierzytelnieniu – przekieruje go pod specjalnie dla niego przygotowany, unikalny adres URL, np. postaci:
<https://example.com/4006A84DAE5B7DD7F13F5A9477...>
- pod którym będzie (z wykorzystaniem mechanizmu *Reverse Proxy*) dostępne dla niego (udostępnione z wykorzystaniem mechanizmu *WebSocket*) okno obsługującej go (i tylko jego) aplikacji stacjonarnej.

Diagram obrazujący opisany powyżej sposób obsługi użytkowników został przedstawiony na rysunku 7.



Rys. 7. Diagram obrazujący sugerowany sposób obsługi użytkowników
(źródło: opracowanie własne)

Warto zauważyć, że opisany sposób działania posiada dodatkową własność, mianowicie nie wykorzystuje *ciasteczek* do utrzymywania sesji użytkownika i identyfikowania żądań kierowanych do serwera przez jego przeglądarkę internetową. Jest to możliwe dlatego, że każdy użytkownik łączy się z innym adresem URL, który jednoznacznie go identyfikuje. Brak użycia *ciasteczek* implikuje niewystępowanie wszystkich (potencjalnych) podatności z nimi związanych w aplikacji internetowej, co jest dodatkową zaletą.

4. Analiza eliminowanych podatności aplikacji internetowych na ataki

W celu potwierdzenia skuteczności proponowanego rozwiązania konieczne jest poddanie analizie potencjalnych podatności (aplikacji internetowych), które mogą zostać wyeliminowane poprzez jego zastosowanie. Jakkolwiek nie istnieje jedna klasyfikacja luk bezpieczeństwa, to uzasadnione wydaje się odwołanie do zestawienia OWASP Top 10, które podaje najczęściej wykorzystywane

podatności podczas ataków na aplikacje internetowe. W celach poglądowych poniżej zostało przytoczone zestawienie z roku 2017 (aktualne w czasie powstawania niniejszego artykułu).

- *A1:2017 – Injection*
- *A2:2017 – Broken Authentication*
- *A3:2017 – Sensitive Data Exposure*
- *A4:2017 – XML External Entities (XXE)*
- *A5:2017 – Broken Access Control*
- *A6:2017 – Security Misconfiguration*
- *A7:2017 – Cross-Site Scripting (XSS)*
- *A8:2017 – Insecure Deserialization*
- *A9:2017 – Using Components with Known Vulnerabilities*
- *A10:2017 – Insufficient Logging & Monitoring.*

Łatwo zauważyć, że niektóre spośród przytoczonych typów podatności nie mogą zostać wyeliminowane przez proponowane rozwiązanie, gdyż są całkowicie niezależne od użytej technologii, języków programowania itp. (mogą wystąpić w każdym przypadku) – zostały one oznaczone na czerwono. Jako przykład można podać *A10:2017 – Insufficient Logging & Monitoring*, gdyż poziom szczegółowości dzienników zdarzeń utrzymywanych po stronie serwera jest ściśle zależny od jego konfiguracji i założeń przyjętych przez jego administratora. Błędy na tym poziomie mogą wystąpić niezależnie od zastosowanego oprogramowania.

Kolorem pomarańczowym zostały oznaczone typy podatności, które jedynie w pewnym (często istotnym) zakresie są eliminowane przez proponowane rozwiązanie.

Powodem występowania większości podatności typu *wstrzyknięcia* (*A1:2017 – Injection*) jest nieprawidłowa obsługa przez aplikację danych pozyskanych od użytkownika – często przyjmowane jest, że są one poprawne i nie są poddawane dodatkowej weryfikacji. Jakkolwiek nie jest możliwe całkowite wyeliminowanie tego typu błędów, to w przypadku aplikacji stacjonarnych sposób przetwarzania wprowadzanych do nich danych zazwyczaj istotnie ogranicza możliwość występowania takich podatności¹². Dodatkowym aspektem jest również to, że w przypadku aplikacji internetowych duża część omawianych

¹² Między innymi ze względu na wykorzystywane biblioteki i środowiska programowe, które zazwyczaj bazują na kodzie binarnym/bajtowym i bezpośrednio nie interpretują kodu źródłowego oraz danych przekazywanych przez użytkowników.

podatności może zostać wykryta jedynie przez tzw. *automaty programowe*, które wykorzystują mechanizm jawnego (w żądaniach HTTP) przekazywania parametrów i ich wartości do serwera w celu modyfikacji przesyłanych danych. Przykładem może być luka bezpieczeństwa znana jako *Time Based SQL Injection*, gdyż jej użycie przez agresora bez wykorzystania specjalistycznych narzędzi jest niemal niemożliwe. Natomiast wykorzystanie obiektu *Canvas* języka HTML5 do interakcji z użytkownikiem powoduje, że wspomniane narzędzia specjalistyczne nie znajdują zastosowania, gdyż dane do aplikacji można przekazać jedynie manualnie, a do serwera są one przesyłane w postaci binarnej (a nie jawnej).

Warto zauważyć, że podobna sytuacja ma miejsce w przypadku podatności typu *A7:2017 – Cross-Site Scripting (XSS)*. Ich występowanie związane jest przede wszystkim z metodą prezentacji danych stosowaną w tradycyjnych aplikacjach internetowych – z wykorzystaniem języka HTML. Dla przykładu – jeśli dane pozyskane od użytkownika zostaną bezpośrednio dołączone do generowanej zawartości witryny, a zawierały one złośliwe sekwencje HTML, to mogą one wpłynąć na wyświetlaną treść. Natomiast w przypadku (istotnej większości) aplikacji stacjonarnych takie zagrożenie w ogóle nie występuje, gdyż ich *Graficzny Interfejs Użytkownika (GUI)* jest tworzony przy wykorzystaniu systemowych/bibliotecznych funkcji graficznych, a nie języka opisu (HTML). Implikuje to brak podatności na ten rodzaj nadużyć.

Zgodnie z definicją podawaną przez *OWASP Top 10*, podatności typu *A2:2017 - Broken Authentication* związane są z ogólnie pojętym logowaniem i zarządzaniem sesją użytkownika. W tym przypadku należy przyjąć, że samo uwierzytelnienie będzie realizowane przez tradycyjną aplikację internetową (która będzie realizowała tylko tę funkcję – patrz punkt 3) i na tym etapie mogą wystąpić ewentualne niedostatki – w zależności od implementacji. Jednakże należy podkreślić, że już po przekierowaniu użytkownika na witrynę z przeznaczoną mu aplikacją nie występują żadne elementy, które mogłyby być podatne na omawiany rodzaj nadużyć (np. *ciasteczka*, co zostało już wskazane w punkcie numer 3).

W przypadku luk bezpieczeństwa typu *A3:2017 – Sensitive Data Exposure* należy zwrócić uwagę, że w zdecydowanej większości przypadków odnoszą się one do udostępniania danych wrażliwych w tradycyjnych aplikacjach internetowych, w których mogą one (dane wrażliwe) zostać zidentyfikowane i pobrane przez agresora przy wykorzystaniu metod typu *Web Scrapping*. Metody te bazują na zautomatyzowanym pozyskiwaniu treści z witryn internetowych oraz na ich maszynowym przetwarzaniu. Jednakże w proponowanej metodzie dane prezentowane są użytkownikowi przy wykorzystaniu obiektu *Canvas* języka HTML5, który nie jest podatny na żadną odmianę metod typu *Web Scrapping*. W tej sytuacji możliwość eksploatacji omawianego typu podatności jest istotnie ograniczona. Pomimo to warto pamiętać, że nie istnieje żadne zabezpieczenie

przed nieumyślnym prezentowaniem użytkownikowi danych wrażliwych, nie dla niego przeznaczonych (np. wyświetlenie danych z konta innego użytkownika bankowości internetowej). Jednakże tego typu sytuacje należą do rzadkości i są zazwyczaj spowodowane istotnymi błędami w logice aplikacji, które nie zostały wykryte na etapie jej testowania przedwdrożeniowego.

Podatności typu *A4:2017 – XML External Entities (XXE)* oraz *A8:2017 – Insecure Deserialization* nie występują, gdyż w rozpatrywanym przypadku nie są do serwera przekazywane dane poddane serializacji, ani zapisane w formie dokumentu XML.

Uwzględniając, że zastosowanie proponowanego rozwiązania oznacza, że w praktyce każdy użytkownik pracuje w ramach wyizolowanego środowiska uruchomionego specjalnie dla niego po stronie serwera, można stwierdzić, że podatności typu *A5:2017 – Broken Access Control* nie znajdują tutaj zastosowania, gdyż z zasady użytkownik ma dostęp jedynie do danych mu przypisanych.

Dla porządku warto zauważyć, że typy podatności eliminowane przez zastosowanie zaproponowanego rozwiązania nie ograniczają się jedynie do tych, które zostały bezpośrednio wymienione w zestawieniu *OWASP Top 10*. Zauważmy, że jedną z częstszych luk bezpieczeństwa jest umożliwianie przeglądarce internetowej zapamiętywania danych wprowadzanych w różnego typu formularzach HTML przez użytkowników¹³. W przypadku zaproponowanego rozwiązania opisany problem w ogóle nie występuje, gdyż wspomniane formularze HTML w ogóle nie są wykorzystywane.

5. Zalety i wady proponowanego rozwiązania

Przedstawiona w artykule koncepcja zwiększenia bezpieczeństwa transakcyjnych aplikacji internetowych bazuje na modyfikacji zarówno ich sposobu działania (wykorzystanie obiektu *Canvas* języka HTML5 do wprowadzania danych), jak i budowy (wykorzystanie dodatkowego mechanizmu *Reverse Proxy*; uruchamianie aplikacji dla każdego użytkownika oddzielnie). Zaproponowane rozwiązanie skutecznie zmniejsza zakres zastosowania dużej części podatności (spotykanych w aplikacjach internetowych) lub w ogóle je eliminuje – co zostało wykazane w punkcie 4 na podstawie luk bezpieczeństwa wymienianych w zestawieniu *OWASP Top 10*. W przypadku aplikacji przetwarzających dane wrażliwe jest to niewątpliwa zaleta, jednakże trzeba mieć

¹³ Działanie to jest często wykorzystywane przez użytkowników ze względu na ich wygodę, jednakże ze względów bezpieczeństwa nie jest ono zalecane i powinno być uznawane za podatność.

także świadomość ograniczeń/wad omawianego podejścia do budowy tego typu programów.

Analizując przedstawioną koncepcję zwiększania bezpieczeństwa aplikacji internetowych, można zauważyć, że jej praktyczne zastosowanie wymaga istotnych zmian w architekturze obecnie funkcjonujących systemów (np. bankowości internetowej). Ich modyfikacja, co oczywiste, wiąże się z wysokimi kosztami, które mogą zostać zaakceptowane tylko przez niektóre instytucje i jedynie w przypadku stwierdzenia, że poczynione inwestycje w przyszłości się zwrócą (co może mieć miejsce po uwzględnieniu nakładów związanych z zabezpieczaniem aplikacji internetowych, wykonywaniem odpowiednich testów penetracyjnych, zakupem specjalizowanych rozwiązań typu *Web Application Firewall* itp.). W związku z powyższym należy przyjąć, że omawiana koncepcja budowy transakcyjnych aplikacji internetowych może znaleźć zastosowanie głównie w nowo projektowanych i implementowanych aplikacjach, co istotnie ogranicza zakres jej zastosowania.

Kolejną wadą, której nie można pominąć, jest kwestia zasobów wykorzystywanych po stronie serwera w przypadku uruchomienia dla każdego użytkownika oddzielnej instancji aplikacji internetowej. Do przeprowadzenia testów empirycznych dla aplikacji posiadającej około 50 widoków wprowadzania danych z 20 aktywnymi polami każdy (co można traktować jako odpowiednik tradycyjnego formularza HTML) konieczne było przydzielenie (zaalokowanie) około 50 MB pamięci RAM. Jest to wielkość istotnie przekraczająca wartości osiąmane przez tradycyjne aplikacje internetowe (często jest to około 5-10 MB dla użytkownika). Należy także zaznaczyć, że w omawianym przypadku pamięć RAM była alokowana na całą sesję użytkownika, zaś przy tradycyjnych aplikacjach internetowych pamięć RAM jest alokowana jedynie podczas generowania odpowiedzi na żądanie użytkownika. W tej sytuacji jednoczesna obsługa wielu sesji użytkowników wymaga zaangażowania sprzętu o istotnie wyższych parametrach niż w przypadku tradycyjnych aplikacji internetowych. Jednakże należy pamiętać, że obecnie efektywność serwerów rośnie, zaś koszty ich nabycia istotnie spadają i często są pomijalne w porównaniu do strat, które instytucja może ponieść w przypadku zakońzonego sukcesem ataku teleinformatycznego na jej zasoby.

6. Podsumowanie

Celem artykułu było przedstawienie nowatorskiej koncepcji zwiększania bezpieczeństwa transakcyjnych aplikacji internetowych, co wydaje się zagadnieniem istotnym ze względu na stale rosnącą liczbę ataków na ten rodzaj oprogramowania (co zostało dokładniej omówione w punkcie pierwszym). W punkcie drugim niniejszego artykułu przytoczono podstawowe informacje na

temat sposobu działania tradycyjnych aplikacji internetowych, zaś w rozdziale trzecim przedstawiono tytułową koncepcję oraz omówiono elementy techniczne, które można wykorzystać do jej wdrożenia. Punkt czwarty poświęcono ogólnej analizie podatności, które mogą zostać wyeliminowane lub których poziom istotności może zostać istotnie zmniejszony poprzez zastosowanie zaproponowanego rozwiązania. Warto nadmienić, że w celach badawczych autorzy przygotowali aplikację internetową zaprojektowaną i zaimplementowaną zgodnie z wytycznymi przedstawionymi w punkcie trzecim. Umożliwiło to praktyczną analizę sposobu działania tego typu oprogramowania oraz wskazanie jego istotnych wad oraz zalet – rezultaty te zostały omówione w punkcie piątym.

Literatura

- [1] BALOCH R., *Ethical Hacking and Penetration Testing Guide*. CRC Press, Boca Raton, 2015.
- [2] CHAUCHAN S., KUMAR N., *Hacking Web Intelligence: Open Source Intelligence and Web Reconnaissance Concepts and Techniques*. Syngress, Waltham, 2015.
- [3] CIAMPA M., *Security Awareness: Applying Practical Security in Your World*. Cengage Learning, Boston, 2015.
- [4] COLLINS M., *Network Security Through Data Analysis*. O'Reilly Media, Sebastopol, 2017.
- [5] PRASAD P., *Mastering Modern Web Penetration Testing*. Packt Publishing, Birmingham, 2016.
- [6] Open Web Application Security Project, *Penetration testing methodologies*. https://www.owasp.org/index.php/Penetration_testing_methodologies
- [7] Open Web Application Security Project, *OWASP Top 10 Application Security Risks – 2017*. https://www.owasp.org/index.php/Top_10-2017_Top_10

The concept for increasing transactional internet applications security

ABSTRACT: The paper presents an innovative concept for increasing the security of crucial Internet applications, for instance ones used in financial operations processing, through non-standard use of HTML5 language mechanisms and server-side processing. Advantages of the proposed approach were discussed, and potential disadvantages of this concept were also pointed out.

KEYWORDS: security, websites, web applications, HTML5, Canvas

Praca wpłynęła do redakcji: 25.05.2018 r.

Zintegrowane badanie i ocena stanu ochrony zasobów informacji

Krzysztof LIDERMAN, Adam E. PATKOWSKI

Instytut Teleinformatyki i Automatyki, Wydział Cybernetyki WAT,
ul. Generała W. Urbanowicza 2, 00-908 Warszawa
krzysztof.liderman@wat.edu.pl, adam.patkowski@wat.edu.pl

STRESZCZENIE: W artykule przedstawiono propozycję zintegrowanego ujęcia zagadnień oceny stanu ochrony informacji w złożonych systemach informacyjnych. Fundamentem tej propozycji jest diagnostyka techniczna oraz bezpieczeństwo informacyjne. Przedstawiono m.in. zagadnienia wykonywania badań dostarczających podstaw do takiej oceny: testów penetracyjnych oraz audytu bezpieczeństwa teleinformatycznego. W ostatnim punkcie opisano krótko metodykę LP-A wykonywania audytu bezpieczeństwa teleinformatycznego integrującą różne typy badań oraz ułatwiającą wykorzystanie różnych, w zależności od potrzeb, wzorców audytowych.

SŁOWA KLUCZOWE: ochrona informacji, bezpieczeństwo informacyjne, diagnostyka, audyt, testy penetracyjne, metodyka LP-A

1. Wprowadzenie

Coraz większa złożoność „świata informacyjnego” sprawia, że osoby odpowiedzialne za szeroko rozumiane przetwarzanie informacji oraz właściciele tych informacji tracą do niego zaufanie i szukają sposobów upewnienia się, że ich informacje są „bezpieczne”. To m.in. powoduje wzrost zleceń na wykonywanie „testów bezpieczeństwa” i „audytów bezpieczeństwa” mających, w rozumieniu zamawiających, dać obiektywną odpowiedź na temat „bezpieczeństwa informacji”, czy też, ujmując problem nieco inaczej – obiektywną ocenę stanu zabezpieczeń informacji. Niestety, podstawowy kłopot polega na tym, że samo pojęcie „bezpieczeństwo” jest rozmyte, a słowa takie jak „testowanie” i „audyt” oznaczają pewne ściśle określone przedsięwzięcia, często znacznie odbiegające od ich popularnego (czytaj także: powierzchownego) rozumienia. Dodatkowo, z powodu łączenia sieci różnych typów (patrz np. [6], [7]), zagadnienia określenia

stanu ochrony informacji przesuwają się z obszaru prostych systemów komputerowych do obszaru złożonych urządzeń i systemów (elektromechanicznych, elektronicznych itp.), w których „element przetwarzania informacji” jest tylko jednym z wielu elementów składających się na takie urządzenia lub systemy.

W niniejszym artykule przedstawiono propozycję zintegrowanego ujęcia zagadnień oceny stanu ochrony informacji w złożonych systemach informacyjnych. Fundamentem tej propozycji jest *diagnostyka techniczna* (patrz punkt 2) oraz perspektywa *bezpieczeństwa informacyjnego*, gdzie bezpieczeństwo informacji jest jego warunkiem koniecznym.

DEFINICJA 1

Bezpieczeństwo informacyjne oznacza uzasadnione zaufanie podmiotu do jakości i dostępności pozyskiwanej oraz wykorzystywanej informacji.

• • • •

Definicja 1 implikuje wykorzystanie kryteriów jakości informacji – są nimi np. relewantność, spójność, dokładność, kompletność, wiarygodność, ale także *bezpieczeństwo*.

DEFINICJA 2

Bezpieczeństwo informacji oznacza uzasadnione zaufanie podmiotu, że nie zostaną poniesione straty wynikające z niepożądanego zmiany, na skutek realizacji zagrożenia, wymaganych wartości istotnych kryteriów jakości informacji.

• • • •

Wymienione w definicji 2 „uzasadnione zaufanie podmiotu” osiąga się zwykle poprzez wykonanie analizy ryzyka i przyjęcie odpowiednich metod postępowania z ryzykiem. Natomiast to jakie kryteria jakości są „istotne” dla poszczególnych kategorii informacji oraz konkretnych uwarunkowań jej przetwarzania, zwykle określają, w przypadku organizacji biznesowej, w której takie informacje są przetwarzane i wykorzystywane, gremia kierownicze tej organizacji (podmiotu) przy pomocy odpowiednich służb. Zwykle kryteriami tymi będą tajność, integralność oraz dostępność informacji. Bardziej szczegółowa dyskusja zagadnień bezpieczeństwa informacyjnego jest zamieszczona w rozdziale 1 w pracy [6].

Ocena to sąd wartościujący, wszelka wypowiedź wyrażająca pozytywne lub negatywne ustosunkowanie się wypowiadającego do kogoś lub czegoś. W dziedzinie techniki, w szczególności w dziedzinie bezpieczeństwa informacyjnego, zasadniczą czynnością procesu oceny jest zbieranie danych stanowiących podstawę opracowania wskaźników, na podstawie których będą wydawane sądy np. o stopniu osiągnięcia założonych celów stawianych przed

zabezpieczeniami. Oceniający ma się zatem wypowiedzieć na temat skuteczności i poprawności spełniania funkcji zabezpieczenia przez obiekt oceniany (czyli potocznie – jakości zabezpieczeń). Ten proces wypracowywania oceny i samą ocenę zwykle nazywa się (niepoprawnie – patrz np. rozdz. 5 w pracy [5]) *pomiarem bezpieczeństwa*.

Na przestrzeni ostatnich kilkunastu lat nastąpiło, pod wpływem m.in. upowszechnienia technologii komunikacyjnych i informatycznych, zwyrodnienie (bo tak to chyba trzeba nazwać) pojęć i przedsięwzięć związanych z określaniem cech jakościowych różnych obiektów, w tym takiego szczególnego obiektu, jak informacja. Wystarczy zajrzeć do będącej w dzisiejszych czasach wyrocznią Wikipedii (polskojęzycznej) żeby stwierdzić, że w dziedzinie techniki *test* i *testowanie* dotyczy tylko oprogramowania – nie bierze się pod uwagę, że są jeszcze inne elementy systemu informacyjnego, które mogą mieć wpływ na jakość (w tym bezpieczeństwo) związanych z nim zasobów informacji, zaś związki pomiędzy takimi pojęciami i przedsięwzięciami, jak *diagnozowanie* i *testowanie* nie istnieją (przynajmniej na poziomie haseł Wikipedii).

2. Diagnostyka techniczna

Dla urządzeń złożonych, jakimi są współczesne systemy cyfrowe, proste metody badań diagnostycznych są często niewystarczające. Dąży się więc do stworzenia i wprowadzenia dokładniejszych metod badania systemów (także w celu automatyzacji procesu ich obsługi) wykorzystujących określony aparat formalny. Jest to obszar dziedziny nauki i techniki nazywany diagnostyką techniczną czy też, w nowszych publikacjach (np. [12]), diagnostyką systemową. *Diagnostyka techniczna* [11] to dziedzina nauki i techniki, zajmująca się opracowaniem oraz wykorzystaniem (czyli teorią i zastosowaniem) metod i środków służących badaniu stanu technicznego systemów. W tym artykule przedmiotem rozważań są złożone systemy cyfrowe (jako podklasa systemów technicznych) i ich dotyczą przedstawione dalej definicje i uwagi.

Badania, na podstawie których jest formułowana diagnoza badanego obiektu, polegają na wykonaniu określonego zestawu sprawdzeń [11]. *Sprawdzeniem* nazywa się ciąg operacji badania wybranej cechy obiektu diagnozowanego, polegających na skontrolowaniu jej wartości i porównaniu ze wzorcem. Podstawowymi czynnościami każdego sprawdzenia są:

- pobudzenie tej części obiektu, od której zależy wartość sprawdzanej cechy;
- odczytanie uzyskanej odpowiedzi (sygnału diagnostycznego);
- porównanie wartości odczytanej z przedziałem wartości dopuszczalnych.

W szczególnym przypadku sprawdzenie może obejmować pojedynczy element (podzespół) obiektu i sprawdzenie takie jest nazywane sprawdzeniem cząstkowym. Częściej jednak sprawdzeniu podlega kilka bądź kilkanaście elementów i w takim przypadku mówi się o torze sprawdzenia. *Torem sprawdzenia* nazywa się podzbiór tych elementów systemu, których stan ma wpływ na wartość badanej cechy diagnostycznej. Proces diagnozowania można zatem zdefiniować następująco:

DEFINICJA 3

Proces diagnozowania to ciąg operacji polegających na wykonaniu określonego zestawu sprawdzeń i analizowaniu uzyskanych wyników w celu rozpoznania stanu, jaki ma miejsce w chwili t_0 zakończenia badania.

• • • •

Diagnozę nazywa się wypowiedź przeprowadzającego badanie o stanie (np. bezpieczeństwa) badanego obiektu. Niech będą dane:

- R – relacje pomiędzy objawami (zestawami wyników sprawdzeń) charakterystycznymi dla poszczególnych stanów a stanami obiektu,
- D – zbiór wyników sprawdzeń (objawów),
- $D_I(t_0)$ – zbiór objawów uzyskany w chwili diagnozowania t_0 ,
- E – zbiór stanów obiektu,
- $\Delta E_I(t_0)$ – diagnoza o stanie E_I dla chwili t_0 .

Można sformułować następującą, tzw. *podstawową regułę diagnostyki* [11]:

DEFINICJA 4

Jeżeli jest znana relacja R oraz uzyskano w procesie diagnozowania konkretny zbiór objawów D_I , to można sformułować diagnozę ΔE_I o stanie obiektu w chwili badania t_0 .

• • • •

Regułę tę symbolicznie można zapisać następująco:

$$R(D_I, E_I) \wedge D_I(t_0) \Rightarrow \Delta E_I(t_0)$$

Czytelnicy obeznani z systemami eksperckimi zapewne zauważą, że reguła ta, w kategoriach terminologicznych systemów eksperckich, jest tzw. *regułą produkcji*.

Warto zwrócić uwagę, że diagnoza sformułowana dla pewnej chwili t_0 jest prawdziwa tylko dla tej chwili t_0 . Ma to odzwierciedlenie w raportach z badań technicznych, głównie testów penetracyjnych – podaje się czasy wykonywania badań oraz aktualność wykorzystywanych baz (eksploatów, poprawek, informacji o podatnościach itp.).

3. Testowanie jako element diagnostyki technicznej

Wprowadzony w poprzednim punkcie termin „sprawdzenie”, używany w diagnostyce technicznej, jest odpowiednikiem używanego w informatyce terminu „test”. W tym artykule rozważane są takie obiekty techniczne, jak systemy informacyjne oraz będące elementami takich systemów obiekty cyfrowe. Mając to na uwadze, podstawowe typy badań wykonywanych w procesie zapewniania jakości (w tym bezpieczeństwa) takich obiektów można wyspecyfikować jako badanie:

1. *Elementów konstrukcyjnych/modułów* – jest to badanie elementów konstrukcyjnych (w przypadku oprogramowania – programu lub modułów wykonywalnych) w celu sprawdzenia, czy nie zawierają one błędów powstałych podczas analizy, projektowania lub wytwarzania. Ponieważ nie wszystkie moduły da się przetestować, korzystając z tzw. pilotów (modułów sterujących przebiegiem testowania) i namiastek (nazywanych też makietami lub zaślepkami) – ich testowanie trzeba wtedy odłożyć do etapu scalania systemu.
2. *Scalania* (ang. *integration test*) – jest to badanie polegające na stopniowym konstruowaniu systemu i wykrywaniu błędów związanych z niepożądanymi zależnościami pomiędzy poszczególnymi elementami konstrukcyjnymi (w przypadku oprogramowania – modułami). Celem jest zbudowanie z oddzielnych, przetestowanych elementów, systemu działającego zgodnie z projektem jego architektury. Scalanie może dotyczyć:
 - modułów w system programowy;
 - oprogramowania ze sprzętem (np. w systemach sterowania).
3. *Zgodności* (ang. *compliance test*) – jest to badanie systemu wykonywane w celu sprawdzenia, czy spełnia on określone wymagania, np. narzucone przez regulacje i standardy państwowe. Pomyślne przejście takich testów może być podstawą do wydania certyfikatu dla systemu.
4. *Systemu* – jest to badanie wykonywane w celu sprawdzenia, czy skonstruowany system spełnia wymagania określone przez użytkowników. Jest to szczególnie przypadek badania zgodności.
5. *Odbiorcze* (ang. *acceptance test*) – jest to badanie systemu lub jego elementu, wykonywane zwykle przez zleceniodawcę (tj. podmiot, który zlecił skonstruowanie tego systemu i jest stroną umowy z wykonawcą), na jego żądanie, po zainstalowaniu w środowisku docelowym, z udziałem wykonawcy w celu sprawdzenia, czy są spełnione wymagania zawarte w umowie.
6. *Weryfikacyjne* (ang. *verification test*) – jest to badanie wykonywane w celu sprawdzenia, czy proces w cyklu rozwojowym systemu spełnia określone

wymagania w danym momencie (etapie) jego opracowywania, tj. czy metody i narzędzia są poprawnie dobrane i stosowane.

Badanie systemu, czyli wytworzonego, złożonego produktu jako całości, wymaga przeprowadzenia pewnych specjalnych testowań, takich jak:

1. *Testowanie wznowień* – polega na powodowaniu lub symulowaniu różnych awarii i sprawdzaniu zdolności systemu do dalszego działania (pożądany jest tzw. stopniowy „upadek” systemu [14]).
2. *Testowanie obciążeniowe* – polega na takim obciążeniu systemu, które spowoduje maksymalne zaangażowanie zasobów systemowych.
3. *Testowanie wrażliwości* – polega na badaniu różnych kombinacji danych wejściowych w celu wykrycia takich, które mogą spowodować niestabilność systemu lub błędy w przetwarzaniu danych lub działaniu.
4. *Testowanie efektywności* – polega na badaniu parametrów operacyjnych (np. czasu odpowiedzi) przy różnych obciążeniach.
5. *Testowanie regresyjne* to powtórne wykonanie niektórych testów w celu upewnienia się, że nowo wprowadzone zmiany nie wywołały niepożądanych skutków ubocznych. Testowanie regresyjne ma szczególne znaczenie na etapie:
 - skalania systemu,
 - eksploatacji (po wykonaniu czynności związanych z tzw. pielęgnacją systemu, np. po zainstalowaniu poprawek lub dołączeniu nowych elementów elektronicznych).

W przypadku testowania systemów sterowania pojawiają się kolejne elementy do przetestowania:

- dotrzymanie wymaganych zależności czasowych w różnych warunkach,
- współdziałanie sprzętu i oprogramowania.

Testowanie wznowień, obciążeniowe, wrażliwości i efektywności są elementami składowymi ogólniejszego *testowania bezpieczeństwa*, które, z grubsza rzecz biorąc, polega na przeprowadzaniu badań systemu pod kątem możliwości nieuprawnionego dostępu do informacji, nieuprawnionej ingerencji w informację oraz możliwości jej uniedostępnienia. Badania takie są zwykle przeprowadzane w ramach audytów, których częścią powinny być wymienione wcześniej badania i testy oraz tzw. *testy penetracyjne* [10] i, w szczególnych przypadkach, badania podatności kodu programu na niepożądane manipulacje.

Testowanie stosuje się w celu wykrycia błędów w testowanym obiekcie. Aby test wykrył błąd, muszą wystąpić trzy elementy:

1. Musi nastąpić *dotarcie do błędu*. Dane wejściowe testu oraz stan testowanego obiektu muszą spowodować uaktywnienie tej części obiektu (np. wykonanie segmentu kodu), w której znajduje się błąd.

2. Musi nastąpić *uwolnienie awarii*. Błąd nie zawsze powoduje złe wyniki, pomimo że element z błędem (obwód elektroniczny, segment kodu itp.) został uaktywniony. Wejście testu oraz stan testowanego obiektu muszą spowodować, że część obiektu, w której tkwi błąd, wyprodukuje błędny wynik.
3. Awaria musi się *ujawnić*. Błędny wynik musi stać się widoczny dla osoby testującej lub automatycznego komparatora.

Użyte terminy „awaria” i „błąd” można zdefiniować następująco¹:

- awaria – jest tak nazywana widoczna niezdolność obiektu testowanego lub jego elementu do wykonania wymaganej funkcji w określonych limitach. Awaria objawia się w postaci niepoprawnego wyjścia (wartości wyjściowych, które nie spełniają wymagań lub specyfikacji), zatrzymania awaryjnego lub niespełnienia wymagań dotyczących określonych nieprzekraczalnych terminów, np. przetwarzania czy przesyłania informacji;
- błąd (ang. *bug*) – jest tak nazywana nieoczekiwana i niezamierzona właściwość programu lub sprzętu, zwłaszcza taka, która powoduje jego wadliwe działanie. Błąd wykryty na etapie eksploatacji obiektu zwykle nazywa się *usterką*. Jako synonimy używane są też terminy: *omylka* oraz *wada*.

Aby stwierdzić, czy obiekt przeszedł pomyślnie test, potrzebna jest tzw. *wyrocznia testowa*. Jest to mechanizm określania (może nim być ekspert), czy rzeczywiste wyniki testu zasługują na akceptację czy odrzucenie. Formalnie rzecz biorąc, potrzebny jest do tego *generator wyników*, który będzie generował wyniki oczekiwane i *komparator*, który będzie porównywał wyniki rzeczywiste z oczekiwanymi. Odrębnym, nie poruszonym w tym artykule zagadnieniem, ale ściśle związanym z testowaniem, jest *zarządzanie błędami* [3].

W tzw. systemach krytycznych ze względu na bezpieczeństwo (ang. *safety-critical systems*) stosowane jest specjalne podejście – w cykl życia systemu wbudowuje się przedsięwzięcia mające na celu uzyskanie odpowiedniego poziomu bezpieczeństwa. Podstawowym przedsięwzięciem jest *analiza bezpieczeństwa* (patrz [14]). W ramach analizy bezpieczeństwa oceniana jest wartość ryzyka związanego z badanym lub projektowanym obiektem i identyfikowane są mechanizmy powstawania szkód w otoczeniu takiego obiektu, w szczególności występowania wypadków z udziałem ludzi. Zdarzenia prowadzące w sposób bezpośredni do wypadku nazywane są *hazardami*, które definiuje się jako sytuacje mogące spowodować śmierć lub obrażenia ludzi.

W projektowaniu systemów tej klasy należy uwzględnić specjalne wymagania, nie występujące przy „zwykłych” systemach informatycznych [13].

¹ W literaturze można spotkać także nieco się różniące objaśnienia podanych dalej terminów.

Do tych specjalnych wymagań, które powinny być zbadane (przetestowane) przed oddaniem systemu do eksploatacji, będą należały wymagania dotyczące:

- przejścia w „stan bezpieczny” w przypadku uszkodzeń i awarii (ang. *fail-safe*). Stan bezpieczny elementu badanego obiektu (systemu) to taki stan, w którym ten element nie stwarza zagrożeń dla obiektu i jego otoczenia, co często osiąga się poprzez ograniczenie funkcji lub dostępności obiektu (systemu),
- detekcji i obsługi błędów (sprzętowych i programowych),
- samotestowania (ang. *selftesting*),
- nadzoru sprzętowego lub programowego nad czasem realizacji funkcji (ang. *watchdog*),
- kontroli uprawnień w trakcie dostępu do zasobów.

Relację elementu obiektu w „stanie bezpiecznym” do pozostałej części obiektu opisują następujące zasady:

- element generuje specjalną sekwencję zdarzeń obserwowalnych (na zewnątrz tego elementu), która pozwala pozostałej części obiektu właściwie zinterpretować aktualny stan elementu;
- nie powstają żadne zdarzenia wewnętrzne, które wyprowadzą ten element ze stanu bezpiecznego – tylko specjalna sekwencja zdarzeń zewnętrznych (obserwowalnych), generowanych przez otoczenie „elementu w stanie bezpiecznym”, może wyprowadzić ten element z tego stanu.

W odróżnieniu od testowania np. systemu programowego, testowanie systemu bezpieczeństwa jest przedsięwzięciem kompleksowym w tym sensie, że obejmuje:

- testowanie poprawności działania zabezpieczeń sprzętowych zintegrowanych z systemem teleinformatycznym,
- testowanie poprawności działania zabezpieczeń programowych zintegrowanych z systemem teleinformatycznym,
- testowanie poprawności integracji ww. zabezpieczeń za pomocą testów skalania i regresyjnych,
- testowanie poprawności działania środków ochrony technicznej i fizycznej,
- sprawdzenie poprawności wdrożenia ochronnych rozwiązań organizacyjnych,
- sprawdzenie odporności personelu na ataki socjotechniczne oraz sprawdzenie współdziałania ww. środków i personelu w celu ochrony informacji przetwarzanej, przechowywanej i przesyłanej w chronionym systemie. To sprawdzenie jest wykonywane za pomocą *testów penetracyjnych* (patrz punkt 4),

- przeprowadzenie testów zgodności ze wzorcem certyfikacyjnym w przypadku wymagania wystawienia certyfikatu bezpieczeństwa lub przeprowadzenie testów odbiorczych.

Zgodnie z ogólnie uznanymi zasadami, przygotowanie do testowania należy zacząć już podczas specyfikowania wymagań na system ochrony, sporządzając m.in. plan testów. Podstawowe informacje niezbędne do zaprojektowania planu testów to ustalenie, według jakiego standardu/normy ma być oceniany docelowy system (czyli jakie „wzorcowe” wymagania powinien spełniać) oraz określenie docelowego poziomu ochrony. Kompleksowe badania systemu bezpieczeństwa przeprowadza się zwykle w ramach audytu tego systemu (zwykle jako *audytu bezpieczeństwa teleinformatycznego*, patrz punkt 5).

4. Testy penetracyjne jako szczególny przypadek testowania

Testy penetracyjne to jeden ze sposobów zbierania informacji w ramach badania bezpieczeństwa systemów, głównie w celu znalezienia podatności. Cechy wyróżniające testy penetracyjne spośród innych rodzajów poszukiwania podatności to [10]:

- potwierdzanie zidentyfikowanych podatności przez precyzyjne wskazanie (często poparte pokazem) skutecznych ataków te podatności wykorzystujących;
- przewaga technik heurystycznych, opartych zwykle na osobistym doświadczeniu realizujących je ekspertów nad ogólnymi, standardowymi technikami „badawczymi” i co się z tym wiąże, na ukrywaniu przez firmy i ekspertów tych sposobów, dających przewagę rynkową nad konkurencją.

Testem penetracyjnym nazywa się także analizę scenariuszy hipotetycznych „ataków”, mającą na celu weryfikację hipotezy o podatności badanego obiektu (lub jego odpowiednika). Analiza taka jest stosowana w przypadkach, gdy rzeczywisty test (np. próba ataku DoS) nie powinien być wykonany.

Należy podkreślić, że test penetracyjny nie jest audytem bezpieczeństwa teleinformatycznego. Podobnie badanie zabezpieczeń teleinformatycznych wyłącznie za pomocą automatycznego skanera to nie jest test penetracyjny; nie jest nim też działanie według wcześniej ustalonej listy kontrolnej (tzw. *checklisty*).

Test penetracyjny to badanie techniczne mające dostarczyć, o ile jest wykonywane w ramach audytu bezpieczeństwa teleinformatycznego (patrz punkt 5), tzw. *dowodów audytowych* do rozstrzygnięcia problemów wskazanych przez audytorów wiodących. Test penetracyjny da odpowiedź na konkretne,

dobrze określone zapytania z zakresu ochrony zasobów informacyjnych; audyt powinien dać odpowiedź na pytanie o stan ochrony zasobów informacyjnych organizacji jako całości, w odniesieniu do ustalonego wzorca.

Celem testu penetracyjnego jest pozyskanie jak największej ilości informacji o testowanym obiekcie i, jeśli będzie to możliwe, również udowodnienie możliwości przełamania zabezpieczeń. Pytanie, na jakie ma dać odpowiedź test penetracyjny to zatem: „jaką informację jest w stanie zdobyć intruz (przy określonych warunkach)?”, a nie „czy da się włamać do systemu?”.

Wynik testu penetracyjnego jest silnie uzależniony od umiejętności i jakości pracy zespołu testującego i od czasu trwania testu. Testy penetracyjne nie gwarantują znalezienia wszystkich podatności, a ich wynik zależy od tzw. frontowych systemów zabezpieczających. Poza tym, czasami trzeba mieć po prostu szczęście i wykonać właściwy test we właściwym czasie.

Ponieważ testy penetracyjne często porównuje się z działaniami intruzów próbujących wykorzystać znane im podatności, warto wspomnieć o zasadniczej różnicy w warunkach działania pentestera i intruza – czas trwania testów penetracyjnych jest ograniczony warunkami umowy pomiędzy zespołem pentesterów a zleceniodawcą, intruz takich ograniczeń nie ma.

Podstawowym opracowaniem wytyczającym trendy w badaniach metodą testów penetracyjnych jest OSSTMM (Open Source Security Testing Methodology Manual). OSSTM jest oficjalną metodyką ISECOM (Institute for Security and Open Methodologies) z zakresu wykonywania testów bezpieczeństwa, udostępnioną na licencji Creative Commons 3.0 Attribution.

Opracowania wspomagające wykonywanie testów penetracyjnych znaleźć można także na stronie fundacji OWASP – Open Web Application Security Project (<http://www.owasp.org>).

Wart polecenia jest także artykuł [10] wiążący metodykę LP-A (patrz dalej) z wykonywaniem testów penetracyjnych oraz opracowania [1] i [2] systematyzujące wykonywanie testów penetracyjnych aplikacji internetowych i prezentujące autorską metodykę PTER ich wykonywania. Szczególnie cenne są zamieszczone tam krótkie, rzeczowe prezentacje metodyk OSSTMM oraz OWASP.

5. Audyt jako szczególny przypadek badania jakości systemu ochrony informacji

Szczególnym rodzajem oceny jest *audyt*. Nazywa się tak postępowanie sprawdzające sposób i/lub wynik wykonania „czegoś” (budowy systemu zabezpieczeń, zarządzania informatyką, wydatkowania funduszy itp.) na

zgodność z określonym wzorcem audytowym (wymaganiami lub standardami albo normami), wykonywane przez stronę niezależną.

Dobrym wzorcem dla audytu bezpieczeństwa jest standard: stanowi rodzaj *checklisty*, daje wektorową miarę „bezpieczeństwa”, zapewnia, że nic, co uznano za istotne nie zostanie pominięte, umożliwia powtarzalność badań. Wzorzec taki jednak słabo uwzględnia specyfikę badanego obiektu oraz nie daje pewności, że zabezpieczenia są szczelne, nie działają przeciw sobie oraz że badanie będzie dostatecznie wnikliwe. Gwoli ścisłości należy nadmienić, że polska Najwyższa Izba Kontroli nieco inaczej definiuje „wzorzec”. Wzorcem według NIK [16] nie jest norma, standard czy specyfikacja wymagań, lecz jednostka sektora prywatnego lub publicznego (krajowa lub zagraniczna – szczegóły patrz rozdziale 5 w pracy [6]).

Wspomniana wcześniej niezależność, w przypadku systemów informacyjnych (i w węższym znaczeniu, teleinformatycznych), musi być zachowana w stosunku do:

- zespołu projektującego lub budującego system informatyczny lub system zabezpieczeń;
- dostawców sprzętu i oprogramowania;
- organizacji podlegającej przeglądowi (w takim sensie, że w skład zespołu audytowego nie powinni wchodzić pracownicy organizacji zlecającej audyt).

Jeżeli ocena ma być rzetelna, to musi być oparta na obiektywnych przesłankach, a nie subiektywnych odczuciach oceniającego. W przypadku audytu tymi obiektywnymi przesłankami są tzw. dowody audytowe. Za [16] dowody kontroli, zgodnie ze Standardami kontroli INTOSAI, nazywa się odpowiednimi, jeżeli są:

- *rzetelne*: wystarczające pod względem ilościowym i właściwe dla uzyskania wyników kontroli/audytu oraz bezstronne i wiarygodne. Wiarygodność dowodów kontroli/audytu zależy od ich charakteru, źródła oraz sposobu ich uzyskania,
- *stosowne*: odnoszące się do celów kontroli/audytu,
- *racjonalne*: koszt zebrania dowodów powinien być proporcjonalny do wyników, które kontroler/audytor zamierza uzyskać.

Te przesłanki (dowody) są uzyskiwane jako wyniki badań, w szczególności badań technicznych – jednym z nich może być testowanie. Zbiór takich informacji (wyników testów) może być podstawą do oceny parametrów jakościowych testowanego obiektu, w tym „bezpieczeństwa”.

Wspomniane pozyskiwanie dowodów audytowych poprzez testy to nic innego jak, używając nieco bardziej technicznych określeń, diagnozowanie, a otrzymane wyniki składają się na diagnozę stawianą na zakończenie audytu.

Przykładem takiej diagnozy może być np. stwierdzenie, że oceniany obiekt spełnia wymagania określone dla wskazanego poziomu EAL standardu Common Criteria (patrz [15]).

Wynikiem oceniania jest raport oceniającego, opisujący zastosowane procedury badawcze wraz z opisem i wynikami testów oraz dokładnie podanym czasem ich wykonania, zastosowanych w celu potwierdzenia lub zaprzeczenia istnienia określonych wad ocenianego systemu teleinformatycznego. Zwykle w raporcie oceniający podaje także sugestie, jakie zidentyfikowane wady i w jakiej kolejności należy usunąć.

Z ogólnej definicji audytu podanej we wstępie punktu wynika, że audyt jest szczególnym przypadkiem oceny. W publikacjach można znaleźć także następujące definicje audytu:

1. „[...] systematyczny, niezależny i udokumentowany proces uzyskiwania dowodu z audytu oraz jego obiektywnej oceny w celu określenia stopnia spełnienia kryteriów audytu”².
2. „[...] niezależny przegląd i badanie zapisów i działań przeprowadzane, by ocenić adekwatność systemowych mechanizmów kontrolnych, zapewnić zgodność z ustalonymi politykami i procedurami operacyjnymi oraz rekomendować konieczne zmiany w mechanizmach kontrolnych, politykach i procedurach”³.
3. „[...] *Audyt wewnętrzny* jest działalnością niezależną i obiektywną, której celem jest przysporzenie wartości i usprawnienie działalności operacyjnej organizacji. Polega na systematycznej i dokonywanej w uporządkowany sposób ocenie procesów: zarządzania ryzykiem, kontroli i ładu organizacyjnego i przyczynia się do poprawy ich działania. Pomaga organizacji osiągnąć cele, dostarczając zapewnienia o skuteczności tych procesów, jak również poprzez doradztwo”⁴.

W Polsce obowiązujące objaśnienia znaczenia terminów związanych z audytem podaje opracowany przez Najwyższą Izbę Kontroli *Glosariusz* [16].

Instytut Audytorów Wewnętrznych (IIA), wyróżnia trzy podstawowe typy audytów: finansowy, zgodności i operacyjny. Wariantem audytu operacyjnego jest *audyt informatyczny*, tj. audyt wykorzystywanych w procesach biznesowych organizacji systemów informatycznych oraz projektów takich systemów. Cele wykonywania audytu informatycznego to przede wszystkim:

² PN-EN ISO 9000:2001 – *Systemy zarządzania jakością. Podstawy i terminologia* (starszą wersję normy przywołano celowo).

³ *National Information Assurance (IA) Glossary*. The Committee on National Security Systems (US), June 2006.

⁴ <https://www.iaa.org.pl/o-nas/definicja-aw> (dostęp: 10.04.2017).

1. Weryfikacja zgodności działania systemów informatycznych z wymogami prawa (np. ustawą *o ochronie danych osobowych* lub *Krajowymi Ramami Interoperacyjności*).
2. Weryfikacja stanu bezpieczeństwa systemów informatycznych (w razie potrzeby także pojedynczych aplikacji) uwzględniająca skutki zidentyfikowanego ryzyka oraz zastosowane procedury kontrolne i ich efektywność.
3. Analiza ryzyka związanego z prowadzeniem projektu informatycznego.

Istotna, chociażby ze względu na koszty przedsięwzięcia i uzyskane wyniki, jest świadomość różnic pomiędzy audytem informatycznym a audytem bezpieczeństwa teleinformatycznego, który jest z kolei wariantem audytu informatycznego. Audyt bezpieczeństwa teleinformatycznego:

- nie dotyczy strony ekonomicznej stosowania środków informatyki w procesach biznesowych;
- nie obejmuje ryzyka związanego z prowadzeniem przez audytowaną organizację projektów z dziedziny informatyki – ocena taka jest zwykle częścią audytu informatycznego;
- nie obejmuje zwykle szczegółowego badania kodu aplikacji i spełnienia wymagań funkcjonalnych i operacyjnych – chyba że jest to wyraźnie zaznaczone w umowie. Standardowo badany jest stopień zaaplikowania poprawek i uaktualnień likwidujących znane podatności.

Na ogólny, praktyczny schemat przeprowadzania audytu bezpieczeństwa teleinformatycznego (zgodny także z opisaną dalej metodyką LP-A), składają się następujące podstawowe czynności:

1. Rozpoznanie (środowiska, stanu posiadania, dokumentacji, procesów) badanej organizacji, np. podmiotu administracji publicznej.
2. Sporządzenie listy audytowej (*checklist*) według wybranego standardu – wzorca audytowego.
3. Wypełnienie listy audytowej z punktu 2 na podstawie ankietowania, wywiadów, wizji lokalnych, kontroli i analizy dokumentów audytowanej organizacji oraz wykonanych testów i badań. Spełnienie zaleceń poszczególnych punktów listy (w miarę potrzeb opatrzone komentarzami) jest kwalifikowane zwykle do jednej z następujących klas: *spełnione*, *nie spełnione*, *spełnione częściowo*, *nie dotyczy*.
4. Badania systemów ochrony fizycznej i technicznej (w tym systemów zasilania w energię elektryczną) oraz sieci i systemów teleinformatycznych eksploatowanych w audytowanej organizacji. Badania te są przeprowadzane przy użyciu wyspecjalizowanych narzędzi i są uzupełniane testami penetracyjnymi.

5. Sporządzenie dokumentacji (raportu) z audytu, obejmującej udokumentowane przedsięwzięcia, wyniki i wnioski dla czynności z punktów 3 i 4. Dokumentacja końcowa musi być podpisana przez audytorów, imiennie gwarantujących rzetelność przeprowadzonej oceny.

Wykraczające poza ramy tego artykułu są pytania o to, jaki jest szczegółowy zakres audytu bezpieczeństwa teleinformatycznego, na ile jego przeprowadzenie może zakłócić normalną pracę organizacji oraz jakiego wsparcia ze strony audytowanej organizacji potrzebują audytorzy (pomocna może być tu metodyka LP-A).

Warto także zauważyć, że w informatyce termin „audyt”, w sensie podanej wcześniej definicji, jest nadużywany. Często w konkretnych umowach z klientem lepiej byłoby używać terminu „ocena”. Termin ten jest znacznie pojemniejszy i łączy się dobrze z terminem „wykonanie” (np. spisu inwentaryzacyjnego). Daje również pole do negocjowania z klientem jego rzeczywistych potrzeb.

Obecnie w środowiskach związanych z audytem informatycznym (głównie dla wariantu audytu – audytu wewnętrznego) toczą się dyskusje nt. przyszłości tego audytu i audytu bezpieczeństwa teleinformatycznego w szczególności oraz próbuje się opisać ujawniające się trendy. Dyskutowane zagadnienia są opisane w rozdziale 5.4 w pracy [6].

6. Metodyka LP-A

Na metodykę LP-A wykonywania audytu bezpieczeństwa teleinformatycznego składają się opisy:

- 1) organizacji, zakresów kompetencji i kwalifikacji zespołu audytowego;
- 2) wyposażenia narzędziowego zespołu audytowego;
- 3) dokumentów niezbędnych do zainicjowania audytu oraz wytwarzanych podczas audytu;
- 4) modelu w postaci diagramów DFD (Data Flow Diagram) opisującego procesy wytwarzania dokumentów i powiązania pomiędzy nimi;
- 5) tzw. „dobrych praktyk”, tj. heurystycznych sposobów postępowania, wypracowanych i sprawdzonych podczas dotychczasowej praktyki audytorskiej twórców metodyki.

Do przedstawienia procesów i przepływu dokumentów podczas audytu zostały wykorzystane elementy metody strukturalnej projektowania systemów informatycznych. Elementy te to przede wszystkim tabele IPO (Input–Process–Output) i diagramy przepływu danych DFD.

Na podstawie metodyki LP-A można:

- 1) ocenić złożoność procesów składających się na audyt;
- 2) prześledzić zależności między dokumentami wytwarzanymi w procesie audytu;
- 3) dobrać skład zespołu audytowego, uwzględniając wymagane zakresy kompetencji (kwalifikacji i uprawnień) członków zespołu;
- 4) ocenić na podstawie zależności między procesami (oraz składu osobowego zespołu) możliwości równoległego prowadzenia zadań audytowych;
- 5) skonstruować harmonogram realizacji audytu (po uwzględnieniu konkretnych warunków wynikających z umowy ze zleceniodawcą);
- 6) oszacować koszty przeprowadzenia audytu (także z uwzględnieniem warunków umowy).

Opublikowana w 2003 roku metodyka [8], [4] była wynikiem doświadczeń i potrzeb związanych z pracami prowadzonymi w tamtym czasie przez jej autorów w zakresie audytów bezpieczeństwa teleinformatycznego. Metodyka została opracowana w celu sprecyzowania i skrócenia dyskusji pomiędzy zleceniodawcą a potencjalnymi wykonawcami na temat „co jest do zrobienia” w pracach, w których zamówienie opiewało na „audyt bezpieczeństwa” lub „sprawdzenie stanu bezpieczeństwa” całości lub części systemu teleinformatycznego. Stwierdzono, że do efektywnego uzgadniania i prowadzenia przedsięwzięć audytowych jest niezbędne opracowanie metodyki, która wspomagałaby osiąganie następujących celów:

1. Opracowanie modelu referencyjnego czynności wykonywanych przez zespół audytowy. Jest on podstawą rozmów pomiędzy wykonawcami i zleceniodawcami o sposobie realizacji audytu z zakresu bezpieczeństwa informacyjnego.
2. Usystematyzowanie przedsięwzięcia audytowego w zakresie wzajemnej kolejności realizowanych czynności audytowych.
3. Usystematyzowanie przedsięwzięcia audytowego w zakresie generowanych w jego trakcie, jak również niezbędnych do jego przeprowadzenia dokumentów.
4. Osiągnięcie komunikatywności i adaptowalności przez wykorzystanie takich mechanizmów formalnych, które:
 - Byłyby zrozumiałe dla osób nieposiadających wiedzy informatycznej (np. nieznających notacji graficznych typu UML i zasad obiektowości). Takie właściwości mają diagramy DFD i tablice IPO, które po krótkich wyjaśnieniach są zrozumiałe nawet dla laików.
 - Byłyby proste w użyciu, w szczególności nie wymuszały stosowania narzędzi skomputeryzowanych.

- Byłyby łatwo skalowalne. Pierwotna wersja metodyki, opublikowana w pracach [4] i [8] została rozwinięta do poziomu 3 diagramów DFD, ponieważ z doświadczeń autorów wynikało, że jest to wystarczający poziom szczegółowości do rozmów z kierownictwem i menedżerami potencjalnego zleceniodawcy audytu. Większa szczegółowość zaciemnia menedżerom (zleceniodawcy) obraz przedsięwzięcia audytowego szczegółami technicznymi. Mniejsza szczegółowość natomiast nie pozwala wykonawcom przedstawić zleceniodawcy potencjalnych problemów wynikających ze skali przedsięwzięcia i zaangażowanych po stronie zleceniodawcy zasobów mających istotny wpływ na czas i koszty prac.

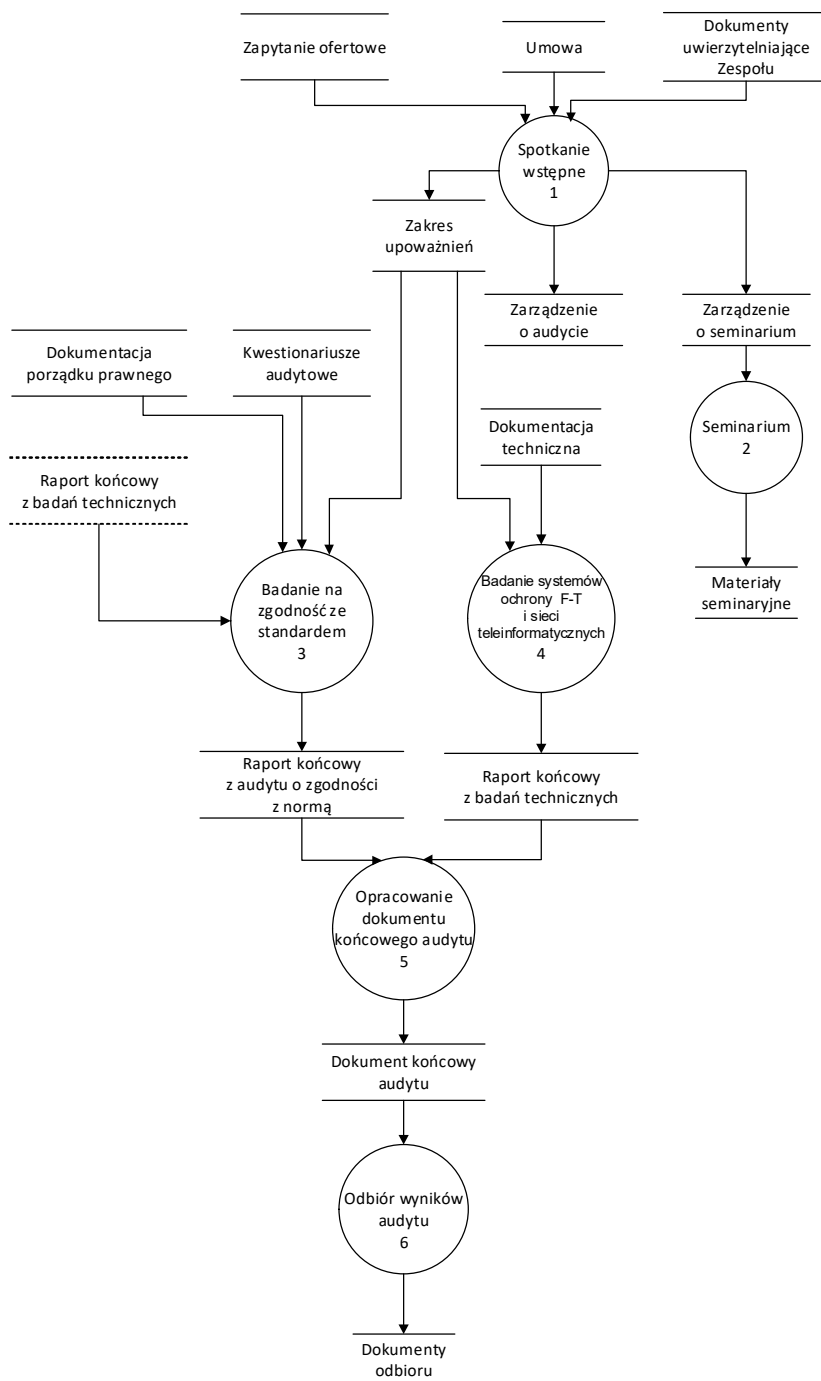
Co prawda audyt, czyli badanie zgodności z pewnym wzorcem, nie odpowiada bezpośrednio na pytanie „czy jest bezpiecznie”, ale rzetelne jego przeprowadzenie zwykle wymaga określenia jakości poszczególnych zabezpieczeń i jakości ich wdrożenia. Ocena ta może być prowadzona w różny sposób, ale autorzy metodyki **uznali techniczne badania bezpieczeństwa za podstawowy element takiej oceny.**

Szczegóły dotyczące rozpowszechnienia metodyki LP-A oraz jej modyfikacji w minionych piętnastu latach zainteresowany Czytelnik znajdzie w pracy [9].

7. Wnioski

W artykule przedstawiono propozycję zintegrowanego podejścia do oceny stanu ochrony informacji w złożonych systemach informacyjnych. Współczesne trendy rozwoju takich systemów, jak „Internet Rzeczy” (IoT), łączenie sieci przemysłowych z sieciami biurowymi, „Big Data” itp. powodują, że ocena stanu ochrony zasobów informacyjnych związanych z takimi systemami, dotąd zwykle ograniczona do oceny stosowanych rozwiązań organizacyjnych i jakości stosowanego oprogramowania, przestaje być wystarczająca.

Podstawowe kryteria jakości „bezpieczeństwa” informacji, takie jak tajność, integralność i dostępność (to ostatnie kryterium w szczególności) zależą bowiem także od zastosowanych i eksploatowanych rozwiązań elektronicznych i elektromechanicznych nośników zasobów informacji i elementów je przetwarzających.



Rys. 1. Diagram DFD procesu audytowego

Elementem integrującym w przedstawionej koncepcji jest metodyka LP-A, której kluczowym elementem są badania techniczne prowadzone w ramach audytu bezpieczeństwa teleinformatycznego, dostarczające dowodów audytowych o osiągnięciu celów stawianych przed zabezpieczeniami zasobów informacji.

Obecnie prowadzone są prace nad wsparciem metodyki LP-A o narzędzie typu „system ekspercki” (patrz definicja 4 i komentarz do niej), wspomagające diagnostykę techniczną prowadzoną na ścieżce technicznej audytu (patrz prawa strona diagramu na rys. 1).

Wejście	- zakresy upoważnień - dokumentacja techniczna sieci i systemów (teleinformatycznych i technicznych) Zleceniodawcy
Nr procesu	4 (*)
Proces	<i>badanie systemów ochrony fizycznej i technicznej oraz systemów i sieci teleinformatycznych Zleceniodawcy (ścieżka techniczna)</i>
Wyjście	raport końcowy z analiz i badań technicznych

Rys. 2. Tabela IPO metodyki LP-A – przykład

Literatura

- [1] ANTCHAK M., ŚWIERCZYŃSKI Z., *Metodyki testowania bezpieczeństwa aplikacji internetowych*. Przegląd Teleinformatyczny, nr 2 (20), 2013, s. 13-38.
- [2] ANTCHAK M., ŚWIERCZYŃSKI Z., *PTER - metodyka testowania bezpieczeństwa aplikacji internetowych*. Przegląd Teleinformatyczny, nr 1-2 (37), 2014, s. 35-68.
- [3] BAYS M.E., *Metodyka wprowadzania oprogramowania na rynek*. WNT, Warszawa, 2002.
- [4] LIDERMAN K., *Podręcznik administratora bezpieczeństwa teleinformatycznego*. MIKOM, Warszawa, 2003.
- [5] LIDERMAN K., *Analiza ryzyka i ochrona informacji w systemach komputerowych*. PWN, Warszawa, 2008.
- [6] LIDERMAN K., *Bezpieczeństwo informacyjne. Nowe wyzwania*. PWN, Warszawa, 2017.

- [7] LIDERMAN K., *Informacyjna ciągłość działania i ataki na sieci różnych typów*, [w:] KOSIŃSKI J. (red.): *Przestępczość teleinformatyczna 2015*. WSPOL, Szcztyno, 2015, s. 47-60.
- [8] PATKOWSKI A.E., LIDERMAN K., *Metodyka przeprowadzania audytu z zakresu bezpieczeństwa teleinformatycznego*. Biuletyn ITA, nr 19, 2003, s. 3-49.
- [9] PATKOWSKI A.E., LIDERMAN K., *Metodyka LP-A – dziesięć lat później*. Przegląd Teleinformatyczny, nr 2 (20), 2013, s. 65-79.
- [10] PATKOWSKI A.E., *Metodyka P-PEN przeprowadzania testów penetracyjnych systemów teleinformatycznych*. Biuletyn ITA, nr 24, 2007, s. 63-96.
- [11] ROZWADOWSKI T., *Diagnostyka techniczna obiektów złożonych*. WAT, Warszawa, 1983.
- [12] ZIELIŃSKI Z., *Podstawy diagnostyki systemowej sieci procesorów o łagodnej degradacji i strukturze hipersześcianu*. WAT, Warszawa, 2013.
- [13] COBIT®5: *Metodyka biznesowa w zakresie nadzoru nad technologiami informatycznymi w przedsiębiorstwie i zarządzania nimi (wersja polska)*. ISACA, 2012.
- [14] IEC 61508: *Functional Safety: Safety-Related Systems*. IEC, 1995.
- [15] www.commoncriteriaportal.org
- [16] NIK: *Glosariusz terminów dotyczących kontroli i audytu w administracji publicznej*. Wyd. 1. Warszawa, lipiec 2005.

Integrated testing and evaluation of the information protection

ABSTRACT: The paper presents a proposal of an integrated approach to the issues of assessing the state of information protection in complex information systems. The foundation of this proposal is technical diagnostics along with information security. Featured, among others issues of performing tests providing the basis for such an assessment: penetration tests and IT security audit. The last chapter of the paper briefly describes the LP-A methodology of performing an IT security audit that integrates various types of research, aiding various audit patterns, depending on the needs.

KEYWORDS: information protection, information security, diagnostics, audit, penetration tests, LP-A methodology.

Praca wpłynęła do redakcji: 10.06.2018 r.

Cyberprzestrzeń jako nowy wymiar aktywności człowieka – analiza pojęciowa obszaru

Maciej MARCZYK

Instytut Teleinformatyki i Automatyki, Wydział Cybernetyki WAT,
ul. Gen. W. Urbanowicza 2, 00-908 Warszawa
maciej.marczyk@wat.edu.pl

STRESZCZENIE: Cyberprzestrzeń jest przestrzenią komunikacyjną tworzoną przez systemy powiązań internetowych. Pozwala jej użytkownikom na komunikację w sieci i nawiązywanie relacji w czasie rzeczywistym. Cyberprzestrzeń jest także środowiskiem wymiany informacji za pomocą sieci i systemów komputerowych. Autor opisuje w artykule przestrzeń cybernetyczną jako wymiar aktywności człowieka, w której wszelkie działania odbiegają charakterem od środowiska fizycznego. Przedstawia zakres pojęciowy i strukturalny tej przestrzeni, która obok środowiska lądowego, morskiego, powietrznego i kosmicznego stała się nowym wymiarem działań człowieka, także tych o charakterze militarnym.

SŁOWA KLUCZOWE: cyberprzestrzeń, sieć teleinformatyczna, zagrożenia

1. Wprowadzenie

Cyberprzestrzeń jest m.in. przestrzenią komunikacyjną tworzoną przez systemy powiązań internetowych. Pozwala jej użytkownikom na komunikację w sieci i nawiązywanie relacji w czasie rzeczywistym. Cyberprzestrzeń jest środowiskiem wymiany informacji za pomocą sieci i systemów komputerowych. Cyberprzestrzeń jest wymiarem aktywności, w której wszelkie działania odbiegają charakterem od środowiska fizycznego. Obok środowiska lądowego, morskiego, powietrznego i kosmicznego, jest to nowy wymiar, w którym można prowadzić działania, w tym działania o charakterze militarnym. Oczywiście jest, że w znacznym stopniu różni się od pozostałych wymiarów, ponieważ:

- jest dziełem człowieka w odróżnieniu od płaszczyzny lądowej, morskiej, powietrznej czy też kosmicznej;
- uczestnicy mają pełną kontrolę nad charakterem tego środowiska;
- nie posiada ograniczeń terytorialnych.

Autorzy piszący o tej przestrzeni najczęściej przedstawiają cztery zasadnicze cechy cyberprzestrzeni, a mianowicie jej anonimowość, aterytorialność, systematyczność i globalny zasięg.

Pojęcie anonimowości odnosi się do możliwości zachowania anonimowości użytkowników cyberprzestrzeni. Należy jednak zaznaczyć, że anonimowość nie wynika z samej istoty cyberprzestrzeni, lecz ze słabości konstrukcji Internetu. Pierwotnie sieć Internet została zaprojektowana dla sektora wojskowego, któremu miała służyć do swobodnej i taniej informacji pomiędzy poszczególnymi placówkami i instytucjami. Wówczas nie myślano o potrzebie ochrony danych i sieci, czy też identyfikacji użytkowników sieci, lecz o prostej konstrukcji, która nie wymagałaby uwierzytelnienia.

Aterytorialność cyberprzestrzeni sprawia, że wszelka aktywność użytkowników w jej obszarze nie nakłada na nich ram/ograniczeń w postaci chociażby granicy geograficznej czy też politycznej. W rzeczywistości oznacza to, że z każdego miejsca na świecie możliwe jest dokonanie połączenia z globalną siecią.

W praktyce ograniczeniem jest jednak:

- przepustowość łącza;
- możliwość technicznego połączenia się z siecią.

Cyberprzestrzeń została ukształtowana przede wszystkim przez proces integracji podstawowych form przekazu i interpretacji informacji, który pozwolił na stworzenie nowych warunków dla funkcjonowania człowieka. Takie zjawisko nazywane może być pojęciem multimedialności, gdzie w procesie wymiany i przekazywania informacji społeczeństwo nie korzysta już jedynie z wymiany werbalnej, lecz z przekazywania informacji za pomocą zdjęć wysyłanych przez Internet czy też filmów umieszczanych na różnych portalach internetowych. Innym istotnym elementem jest ciągła konwergencja systemów i usług umożliwiających komunikację czy też przesyłanie danych. Obecnie coraz bardziej zauważalne staje się zjawisko integracji technicznej, czyli stworzenie zintegrowanej platformy teleinformatycznej służącej do międzynarodowej komunikacji społecznej.

2. Wybrane naukowe spojrzenia na cyberprzestrzeń

Dotychczasowe rozważania pozwalają na przedstawienie cyberprzestrzeni w różnych ujęciach. Zdaniem autora, w zależności od dziedziny nauki, definicje

cyberprzestrzeni będą inaczej obrazowane, a przede wszystkim w różnym znaczeniu interpretowane. Inne spojrzenie na cyberprzestrzeń będą przedstawiać nauki techniczne, jeszcze inne nauki społeczne, socjologiczne czy też psychologiczne.

W obszarze nauk społecznych cyberprzestrzeń jest płaszczyzną służącą do komunikacji międzyludzkiej, jak również nowym wymiarem rozrywki dla społeczeństwa. Józef Bednarek wyróżnił pięć podstawowych funkcji cyberprzestrzeni, do których zaliczamy [1]:

- funkcję informacyjną;
- funkcję ludyczną (rozrywkową);
- funkcję stymulującą;
- funkcję wzorotwórczą;
- funkcję interpersonalną.

Funkcja informacyjna cyberprzestrzeni określa, iż jest ona przestrzenią ciągłej wymiany informacji. Pozwala ona na błyskawiczne uzyskanie różnorodnych informacji z niemalże każdego aspektu życia człowieka czy też nauk. Istota funkcji ludycznej pozwala zaobserwować, iż cyberprzestrzeń pozwala na zaspokojenie potrzeb człowieka związanych z rozrywką czy oderwaniem się od codziennych zajęć. Możliwe jest to dzięki coraz większemu dostępowi do infrastruktury teleinformatycznej. Funkcja stymulująca wyraża się w inspiracji odbiorców do aktywnego odbioru treści znajdujących się w cyberprzestrzeni. Kolejna funkcja odnosi się do zjawisk tworzenia nowych standardów, ideałów czy też stylów życia, które obserwujemy wraz z rozwojem cywilizacyjnym. Coraz częściej spotykane są próby definiowania pojęcia cyberspołeczności czy też wirtualnej rzeczywistości.

Socjologiczne podejście do cyberprzestrzeni opisuje ją jako przestrzeń, gdzie następuje mieszanie się kultur, gromadząc przy tym różne idee oraz informacje z całego świata. Pojęcie cyberprzestrzeni definiowane jest jako *Nowa Wieża Babel* czy też *Cyber Termopile*, gdzie dochodzi do zacierania się różnic pomiędzy kulturami, jak również jest obszarem ciągłych starć i konfliktów [11].

Cybernetyczne podejście do pojęcia cyberprzestrzeni definiuje ją jako przestrzeń otwartej komunikacji, gdzie sprzężenie zwrotne informacji pozwala na regulację systemów, zachodzenie relacji między nimi oraz dynamiczny rozwój i wytyczanie nowych szlaków [6].

Podejście psychologiczne to podejście, gdzie pojęcie cyberprzestrzeni utożsamiane jest z wirtualną rzeczywistością, jako przestrzeń zapośredniczona w kontekście technologii informacyjnej, w której obecne są wirtualne byty, niedostępne dla człowieka w ich cyfrowym, abstrakcyjnym wymiarze, zobrazowane poprzez interfejs w wymiarze dostrzeganym wszelkimi zmysłami poprzez dźwięk lub też obraz [6].

Powyższe rozważania dotyczące cyberprzestrzeni definiują ją jako z informatyzowaną przestrzeń wymiany informacji.

Największe znaczenie w procesie wymiany informacji posiada Internet, czyli termin oznaczający ogólnosiwiatową sieć komputerową służącą wymianie informacji za pomocą znormalizowanych protokołów komunikacyjnych¹. Za pomocą tej globalnej sieci użytkownicy z całego świata są w stanie komunikować się między sobą.

Dużego znaczenia w tym procesie nabiera infrastruktura teleinformatyczna, czyli zespół środków teleinformatycznych, takich jak sieci czy systemy teleinformatyczne. W literaturze przedmiotu wyróżnić można liczne synonimy pojęcia Internet, takie jak globalna sieć komputerowa, ekstranet, sieć ogólnosiwiatowa, jak również w języku potocznym spotykane określenie net. Pojęcie Internetu (z ang. Inter-Network) w dosłownym tłumaczeniu to między sieć.

Można wyróżnić cztery zasadnicze elementy, które charakteryzują tę globalną sieć [3]:

- nieograniczone zasoby informacji i danych dostępne i znajdujące się w Internecie;
- interaktywność oznacza, iż w czasie rzeczywistym użytkownicy mogą się ze sobą komunikować;
- powszechność definiowana poprzez coraz łatwiejszy dostęp, jak również łatwość i coraz większe jej znaczenie w niemal każdym aspekcie życia człowieka;
- demokracja, w myśl której globalna sieć określana jest najbardziej demokratycznym medium z możliwością całodobowego dostępu.

Zdaniem autora ważny jest również aspekt ekonomiczny, który definiuje globalną sieć jako jeden z najtańszych środków wieloaspektowych związanych z przetwarzaniem informacji czy źródeł rozrywki w stosunku do kosztów ponoszonych w wyniku dostępu do niego. Istotna jest również cecha powszechności, gdyż obecnie coraz większa część społeczeństwa nie wyobraża sobie życia bez szybkiego dostępu do Internetu.

3. Wybrane pojęcia w obszarze cyberprzestrzeni

Z pojęciem cyberprzestrzeni wiąże się kilka zasadniczych elementów definiujących ją jako pewien spójny wymiar. Wśród tych elementów

¹ <http://www.oxforddictionaries.com/definition/english/Internet>

wyróżniamy: *systemy i sieci teleinformatyczne, dane i informację oraz użytkowników.*

System to wyodrębniony zbiór elementów materialnych lub abstrakcyjnych, wzajemnie powiązanych, jako całość z określonego punktu widzenia, mający przy tym takie właściwości, których nie posiadają jego elementy [8]. W dziedzinie technologii informacyjnej system będzie określany zbiorem powiązanych elementów służących do przetwarzania danych przez określone środki informatyczne.

Pojęcie **systemu teleinformatycznego** oznacza zespół współpracujących ze sobą urządzeń informatycznych i oprogramowania, zapewniający przetwarzanie i przechowywanie, a także wysyłanie i odbieranie danych poprzez sieci telekomunikacyjne za pomocą właściwego dla danego rodzaju sieci telekomunikacyjnego urządzenia końcowego [15]. Pojęcie systemu teleinformatycznego wiąże się więc ze zbiorem określonych urządzeń informatycznych, które służą do przetwarzania informacji za pomocą określonego medium transmisyjnego do urządzeń końcowych. Pojęcie medium transmisyjnego oznacza nośnik używany do transmisji sygnałów w telekomunikacji, umożliwiający rozchodzenie się fal akustycznych, elektrycznych, radiowych i świetlnych.

W środowisku cybernetycznym możemy wyróżnić sieci: teleinformatyczne, komputerowe lub telekomunikacyjne.

Analiza aktów prawnych wykazała, że **sieci teleinformatyczne to sieci pozwalające na wysyłanie i odbieranie danych pomiędzy systemami informatycznymi pełniącymi rolę urządzeń końcowych** [16]. Józef Janczak określa natomiast sieć teleinformatyczną jako sieć pozwalającą na odbieranie oraz wysyłanie danych pomiędzy systemami teleinformatycznymi określanymi jako urządzenia końcowe [4]. Literatura przedmiotu pozwala określić sieć teleinformatyczną jako zbiór urządzeń, takich jak router, switch czy modem, wykorzystywanych do przekazywania informacji wysyłanych ze stacji roboczych [9]. Sieć teleinformatyczna może być więc infrastrukturą służącą zarówno do nadawania, transmisji, jak i odbioru informacji przedstawionej w postaci cyfrowej². Zdaniem autora sieci teleinformatyczne to pewna architektura powiązań organizacyjnych i technicznych wykorzystywanych do połączenia systemów teleinformatycznych. Wykorzystuje ona specjalistyczne elementy teleinformatyczne służące do przekazywania informacji pomiędzy poszczególnymi urządzeniami informatycznymi.

Sieci komputerowe, w literaturze przedmiotu określane głównie jako sieci informatyczne, to sieci służące do wymiany informacji pomiędzy osobami funkcyjnymi stanowisk dowodzenia wyposażonymi w komputery lub inne

² http://epodlaskie.wrotapodlasia.pl/pl/dzialania/konferencja_informacyjna.htm?m=dload&debug=off&id=53

urządzenia informatyczne. Sieci komputerowe służą przede wszystkim ułatwieniu komunikacji pomiędzy poszczególnymi jej użytkownikami zarówno w sektorze prywatnym, jak i w strukturach zhierarchizowanych. Sieci te pozwalają również w szybki sposób uzyskać dostęp zarówno do zasobów danej organizacji, jak i globalnych informacji za pośrednictwem Internetu. Sieci komputerowe pozwalają wykonywać pracę zdalnie bez fizycznego kontaktu z danym urządzeniem, jak również pozwalają na udostępnienie zasobów sprzętowych podłączonych w środowisku sieciowym. Przykładem może być dostęp zdalny do urządzeń peryferyjnych, takich jak skanery, drukarki sieciowe czy pamięci masowe oraz dyski wirtualne. W sieciach komputerowych wyróżnia się również dwa zasadnicze komponenty, takie jak elementy aktywne i elementy pasywne sieci. Do elementów aktywnych zaliczyć można karty sieciowe, wzmacniacze, routery punkty dostępowe, jak również koncentratory bądź przełączniki. Do elementów pasywnych zaliczamy kanały kablowe, pomieszczenia techniczne, maszty transmisyjne itp. Warto zaznaczyć, iż sieci komputerowe mogą ograniczać się do jednego bądź kilku budynków, ale mogą też pokrywać obszary miasta, kraju, a nawet kontynentów.

Jednym z podstawowych kryteriów podziału sieci komputerowych jest obszar, stąd wyróżnić możemy³:

- Sieci lokalne (LAN) zwane lokalnymi sieciami komputerowymi obejmującymi zazwyczaj małe obszary, takie jak dom, biuro czy budynek.
- Sieci miejskie (MAN) obejmujące swoim obszarem zazwyczaj miasto łączy, kilka lokalnych sieci komputerowych.
- Sieci rozległe (WAN) obejmujące swym obszarem całe państwa i kontynenty. Przykładem takiej sieci może być Internet.

Kolejnym kryterium jest funkcja, stąd wyróżnić możemy:

- WLAN (ang. Wireless Local Area Networks), czyli bezprzewodową sieć lokalną działającą na małym obszarze, w której urządzenia połączone są ze sobą przy użyciu komunikacji bezprzewodowej.
- SAN (ang. Storage Area Networks) zwane siecią pamięci masowej, która zapewnia systemom informatycznym dostęp do zasobów pamięci masowej. Obecnie używana jedynie w dużych korporacjach ze względu na złożoną strukturę i ogromne koszty zarówno implementacji, jak i utrzymania.
- CAN (ang. Controller Area Networks) jest to szeregową magistrala komunikacyjna odporna na zakłócenia elektromagnetyczne dzięki zastosowaniu różnicowej transmisji bitów. Stosowana pośrednio w przemyśle, np. w systemach samochodowych, gdzie istotna jest właśnie transmisja danych, a niekoniecznie jej szybkość przesyłania. W tym rodzaju

³ <http://wireless-network-support.blogspot.com/2009/08/what-is-lan-wlan-wan-man-san-can-pan.html>

sieci maksymalna prędkość przesyłania danych to 1 Mb/s w odległości do 40 metrów. Prędkość przesyłania danych spada wraz ze wzrostem odległości.

- PAN (ang. Personal Area Networks) – rodzaj sieci komputerowej używanej do transmisji danych poprzez różne urządzenia użytkownika. Obecnie taką sieć można stworzyć za pomocą połączenia komputera, telefonu komórkowego (bądź smartfona), telewizora, tabletu czy innych urządzeń użytkownika w jedną sieć wymiany danych.

Sieci telekomunikacyjne według Ustawy *Prawo telekomunikacyjne* to systemy transmisyjne oraz urządzenia komutacyjne lub przekierowujące, a także inne zasoby, w tym nieaktywne elementy sieci, które umożliwiają nadawanie, odbiór lub transmisję sygnałów za pomocą przewodów, fal radiowych, optycznych lub innych środków wykorzystujących energię elektromagnetyczną, niezależnie od ich rodzaju [17]. Jest to zespół urządzeń telekomunikacyjnych znajdujących się na danym obszarze, przeznaczony do świadczenia usług telekomunikacyjnych.

W literaturze przedmiotu pomiędzy sieciami telekomunikacyjnymi i teleinformatycznymi można zaobserwować pewne prawidłowości związane z zależnościami funkcjonalnymi, które pozwalają porównać poszczególne przeznaczenia sieci [4]:

- zarządzanie transmisją informacji;
- wykorzystanie systemów teletransmisyjnych;
- generowanie sygnałów z urządzeń do systemów transmisyjnych za pomocą określonych częstotliwości;
- synchronizacja pomiędzy odbiornikiem i nadajnikiem;
- adresowanie pakietów na kierunku odbiorca – nadawca;
- wykrywanie i korekta błędów danych;
- odpowiedni dobór drogi przepływu pakietu;
- retransmisja z powodu błędów sieci i awarii;
- formatowanie danych zawartych w wiadomości oraz ustalenie formatu wymiany wiadomości pomiędzy poszczególnymi urządzeniami końcowymi (stacjami roboczymi);
- integracja zarządzania siecią oraz bezpieczeństwa sieci.

Kolejnym przytoczonym przez autora elementem są *dane i informacje*. **Dane** to informacje przedstawione w postaci umożliwiającej ich przetwarzanie za pomocą programów komputerowych lub będące wynikiem tego przetwarzania⁴. Dane reprezentują fakty i są przesyłane do świadomości

⁴ <http://encyklopedia.pwn.pl/haslo/3890542/dane.html>,

odbiorcy w postaci komunikatu. Danymi można nazwać również pewne zbiory faktów, zjawisk, zdarzeń, cech przedmiotów lub obiektów, które są dostępne w bazach danych lub uzyskuje się je dzięki sensorom [10]. W systemach i sieciach teleinformatycznych wyróżnione zostało również pojęcie *danych osobowych*, czyli wszelkich informacji dotyczących zidentyfikowania lub możliwych do zidentyfikowania osoby fizycznej [18].

W sieciach teleinformatycznych wyróżniamy również dwa zasadnicze typy danych w nich przetwarzanych [18]:

- Dane lokalizacyjne, czyli te dane, które przedstawiają położenie geograficzne danego podmiotu korzystającego z usług telekomunikacyjnych. Przykładem takich danych może być długość i szerokość geograficzna, wysokość nad poziomem morza. Dzięki tego typu danym możliwe jest określenie kierunku poruszania się danego obiektu, jak również określenie jego dokładnej pozycji czy prędkości, z jaką się porusza.
- Dane telekomunikacyjne odnoszące się bezpośrednio do danych usług sieciowych, z jakich korzysta użytkownik. Przykładami takich danych może być np. rodzaj sieci źródłowej i końcowej, wykorzystywany przez użytkownika danej sieci protokół, ilość wysyłanych i odbieranych danych w trakcie połączenia, jak również czas sesji, czy wybór trasy pakietów.

Pojęcie **informacji** natomiast jest niewątpliwie trudniejsze do zdefiniowania, bowiem różne dziedziny nauki definiują to pojęcie w zależności od charakteru, czy też sposobu użycia. Jest to także zbyt szeroki obszar na jego opisanie w tym artykule. Najogólniej można przyjąć, że jest to określona treść przedstawiona za pomocą określonego języka lub kodu przez nadawcę do odbiorcy [14]. Połączenie informacji i danych przedstawił Jerzy Kisielnicki pisząc, że informacja to przekazywanie różnorodności, a jej znakową postacią są dane [5].

Pojęcia danych i informacji w języku potocznym często są używane zamiennie, a nieprecyzyjne określanie tych pojęć prowadzi do niewłaściwego zrozumienia samej istoty przekazywania informacji.

4. Wybrane pojęcia zagrożeń cyberprzestrzeni

Najogólniej **zagrożenie** jest interpretowane jako sytuacja lub stan, w którym komuś zagrażają lub w którym ktoś czuje się zagrożony [2]. W podejściu zarządzania kryzysowego zagrożenie będzie rozumiane poprzez pewne następstwo zdarzenia, które wywiera negatywny wpływ na funkcjonowanie politycznych i gospodarczych struktur państwa, czy też na warunki bytowania ludności oraz stan środowiska naturalnego [7]. Zagrożenie

jest więc pewnym zdarzeniem, które powoduje zmniejszenie poziomu bezpieczeństwa danego podmiotu. Zdaniem autora pojęcie „zagrożenie” ma charakter interdyscyplinarny i w zależności od interpretowanej nauki, pojęcie może nabierać innego znaczenia. Zagrożenie można utożsamiać z pewnym stanem, jak również następstwem straty danej cechy czy też przedmiotu. Zagrożenie wiąże się więc z pewnym stanem, który powoduje powstanie pewnej niebezpiecznej sytuacji dla danego podmiotu bądź otoczenia. Zdaniem autora kluczowe w terminologii jest również określenie, iż zagrożenie cechuje się pewnym prawdopodobieństwem wystąpienia. Mimo że dany podmiot jest zagrożony, to istotne jest określenie szansy wystąpienia danego zjawiska odbierającego poczucie bezpieczeństwa.

Zdaniem autora z pojęciem zagrożenia wiąże się również termin **bezpieczeństwa**, czyli stan, który daje poczucie pewności i gwarancję jego zachowania oraz szansę na doskonalenie. Pojęcie bezpieczeństwa to jednak, podobnie jak znaczenie informacji, zbyt rozległy temat badawczy, nie mieszczący się w ramach tematyki tego artykułu.

Analiza literatury przedmiotu wykazała jednak, że oprócz pojęcia zagrożeń cyberprzestrzeni, spotykane są również definicje cyberzagrożeń, zagrożeń w środowisku cybernetycznym, czy też zagrożeń obszaru cyberprzestrzeni [12]. Zagrożenia dość często interpretowane są jako możliwości negatywnego wpływu na działanie podmiotów w cyberprzestrzeni. W przypadku zagrożeń w cyberprzestrzeni odnosząc się do aspektu przetwarzanych informacji, można określić, że są nimi zdarzenia losowe bądź działania celowe mogące spowodować przesłanki do legalnego lub nielegalnego ujawnienia informacji, jej modyfikacji, zniszczenia czy ewentualnej kradzieży.

Zdaniem autora w procesie identyfikacji zagrożeń kluczowe jest zdefiniowanie środowiska cybernetycznego. Pojęcie to określane jest jako zespół wszelkich elementów i czynników będących w ścisłej współzależności, który wpływa na procesy informacyjne danego układu poprzez umacnianie stanów pożądanых i przeciwdziałanie owym stanom niepożądanym [13]. Elementami tego środowiska są wszelkie stacje robocze, terminale, koncentratory, serwery, jak również w odniesieniu do SZ RP wszelkie radiostacje czy też radiolinie.

Ze względu na ich źródło pochodzenia wyróżnia się zagrożenia wewnętrzne i zagrożenia zewnętrzne.

Zagrożenia wewnętrzne dotyczyć będą wszelkich działań odnoszących się pośrednio lub bezpośrednio do elementów cyberprzestrzeni, na które oddziałuje podmiot wewnętrzny danego otoczenia. Zagrożenia zewnętrzne natomiast to takie zjawiska, które powstają na skutek działalności zewnętrznych podmiotów na dane środowisko. W przypadku zagrożeń wewnętrznych odnoszących się do środowiska sieci i systemów komputerowych będą to wszelkie działania

zachodzące wewnątrz danego systemu teleinformatycznego. Wśród nich wyróżnić można:

- nieuprawniony dostęp do zasobów sieci;
- kradzież informacji i danych przez użytkowników wewnętrznych;
- niewłaściwe wykorzystanie aplikacji działających w systemie;
- świadome wykorzystywanie luk systemów i oprogramowania;
- awarie sprzętowe, łączy, czy też oprogramowania,
- inne działania związane z wypadkami bądź spowodowane czynnikiem ludzkim.

Nieuprawniony dostęp do zasobów sieci oznacza, iż użytkownik do tego nieupoważniony jest w stanie pozyskiwać informacje, do których nie powinien mieć dostępu. Nieautoryzowany dostęp do informacji można uzyskać w trzech obszarach. Pierwszym jest obszar centrali, gdy użytkownik nieuprawniony uzyskuje dostęp w głównych punktach dostępowych, takich jak siedziba główna danej firmy czy też organy bądź obiekty zarządzające na szczeblu strategicznym. Przykładem może być siedziba główna banku czy też departament IT zajmujący się zarządzaniem sieciami i systemami teleinformatycznymi. Drugim obszarem jest określone stanowisko robocze, gdzie użytkownik za pomocą danego urządzenia jest w stanie uzyskać dostęp do poufnych informacji. Ostatnim jest obszar jest związany z liniami transmisyjnymi, gdzie napastnik uzyskuje dostęp do danych poprzez włamanie się w procesie ich transmisji. Pojęcie nieuprawnionego dostępu do zasobów sieci można nazywać również hackingiem komputerowym. Hacking komputerowy to pojęcie związane z działalnością użytkownika, który bez uprawnienia uzyskuje dostęp do informacji dla niego nieprzeznaczonych, podłączając się do sieci telekomunikacyjnej lub przełamując albo omijając elektroniczne, magnetyczne, informatyczne lub inne szczególne jej zabezpieczenia⁵.

Kradzież informacji i danych przez użytkowników wewnętrznych oznacza działalność mającą na celu pozyskiwanie danych, a w szczególności danych poufnych. Zdaniem autora kluczowe jest określenie podmiotów realizujących powyższe działania.

Wyróżniamy trzy zasadnicze podmioty pozyskujące informację i dane poufne⁶: Hacker, Cracker i Phreaker.

Pojęcie hackera określa osobę uzyskującą dostęp do nieautoryzowanych informacji. Wśród typów hackerów wyróżnia się Black Hat, czyli osoby działające na pograniczu prawa, aby móc włamywać się do systemów i sieci

⁵ http://www.cyberprzestepczosc.info/hacking_komputerowy.html

⁶ <http://gadzetomania.pl/11106,haker-cracker-phreaker-czym-roznia-sie-od-siebie-sieciowi-przestepcy>

komputerowych bądź omijać zabezpieczenia, korzystając z luk w ochronie. Celem jest pozyskiwanie informacji bądź chęć ujawnienia dostępu do owych danych, aby pokazać swoje możliwości bądź podnieść własną wartość jako osoby posiadającej specjalistyczną wiedzę. Istnieją jeszcze hakerzy zaliczający się do grupy White Hat, czyli osób, które z założenia unikają wyrządzenia poważnych szkód. Mają raczej na celu pokazanie danym instytucjom czy firmom luk w ich zabezpieczeniach oraz hakerzy, którzy przyjmują po części metody działania obu wyżej wymienionych grup, nazywani Grey Hat.

Cracker to osoba świadomie łamiąca zabezpieczenia, ale w celu uzyskania własnych korzyści. Crackerzy często działając, myślą o szkodzie dla danego systemu czy sieci teleinformatycznej bądź świadomej kradzieży informacji poufnych w celach finansowych.

Ostatnią grupą są tzw. Phreakerzy, czyli osoby łamiące zabezpieczenia sieci telefonicznych w celu uzyskania darmowych połączeń bez konieczności wpinania się w linię abonenta. Jest to jednak grupa, obecnie coraz mniej spotykana, gdyż obecnie usługi telekomunikacyjne są coraz powszechniejsze i coraz tańsze.

Niewłaściwe wykorzystanie aplikacji działających w systemie oznacza, iż użytkownicy mogą nieświadomie wykorzystywać oprogramowanie w ten sposób, że narażają funkcjonowanie całego systemu, otwierając drogę do ujawnienia informacji nie tylko swoich, lecz również informacji, które mogą być utajnione. Niewłaściwe korzystanie z aplikacji może prowadzić do⁷:

- dostępu osób trzecich do informacji prywatnych, firmowych czy innych poufnych danych;
- ujawnienia haseł do systemów czy innych aplikacji;
- umożliwienia monitorowania urządzeń poprzez inne aplikacje działające w tle systemów operacyjnych;
- umożliwienia śledzenia położenia danego użytkownika;
- umożliwienia uszkodzenia urządzeń teleinformatycznych;
- wykorzystania prawdziwej tożsamości użytkownika do innych celów.

Świadome wykorzystywanie luk systemów i oprogramowania może prowadzić do innych działań, wliczając w to kradzież informacji i ujawnienie informacji poufnych. Problematyka wykorzystywania luk w systemach i oprogramowaniu dotyczy w głównej mierze źle zaprojektowanych kodów aplikacji, porzuconej cyfrowej własności oraz innych błędów popełnianych przez użytkowników/pracowników⁸.

⁷ <http://bitdefender.marken.com.pl/2014/clueful>

⁸ <http://www.cisco.com/web/PL/prasa/news/2014/20140807.html>

Klasyfikacja zagrożeń cyberprzestrzeni w literaturze przedmiotu nie pozwala precyzyjnie określić tych, które stanowią największe zagrożenie dla funkcjonowania i bezpieczeństwa państwa. Możliwe jest jednak określenie tych, które występują najczęściej.

W opinii autora czynnikami, które powodują przeniesienie obecnych działań przestępczych, jak również militarnych, w sferę cyberprzestrzeni są:

- Niski koszt utrzymania jednostek militarnych (wojsko, sprzęt, amunicja itp.) w stosunku do małych kosztów utrzymania kilku osób mogących sparaliżować duży system, armię czy państwo (jednostki informatyków/hackerów).
- Trudne do udowodnienia przestępcze działanie w cyberprzestrzeni, rzadko udaje się w pełni obarczyć winą określony podmiot.
- Nieadekwatne rodzaje kar i egzekwowanie prawa, mniejsze konsekwencje prawne dotyczą kradzieży informacji, danych, jak również środków finansowych przy użyciu cyberataków aniżeli w rzeczywistym realnym świecie.
- Możliwość ataku w cyberprzestrzeni z niemal każdego miejsca na świecie powoduje, iż nie jest potrzebny bliski kontakt w przypadku prowadzenia działań.

Podsumowując, cyberprzestrzeń stała się nowym środowiskiem, które pozwala w dużym stopniu rozwijać się patologiom związanym z działaniami grup przestępczych, terrorystów czy też oszustów. Na społeczności międzynarodowej spoczywa ogromna odpowiedzialność związana z tworzeniem procedur, mechanizmów i standardów poprawiających bezpieczeństwo w globalnej sieci komputerowej.

5. Zakończenie

W ramach przeprowadzonej analizy pojęciowej można stwierdzić, że cyberprzestrzeń jest nową przestrzenią komunikacyjną tworzoną przez systemy powiązań internetowych i potwierdzić założenie wstępne, że pozwala jej użytkownikom na łatwiejszą komunikację w sieci i nawiązywanie relacji w czasie rzeczywistym. Jest zatem nowym wymiarem aktywności człowieka, w której wszelkie działania odbiegają charakterem od środowiska fizycznego.

Zdaniem autora działania podejmowane na rzecz ochrony cyberprzestrzeni nigdy nie będą jednak wystarczające, aby jakkolwiek podmiot mógł w pełni czuć się bezpieczny. Nieograniczona ilość zagrożeń, jakie występują, przy zbyt małej ilości rozwiązań wpływających na jej bezpieczeństwo, nie napawa optymizmem. Jednak nadchodzące lata mogą

w znacznym stopniu ukierunkować działania na korzyść zwykłych użytkowników sieci i bezpiecznej komunikacji między nimi.

Literatura

- [1] BEDNAREK J., *Cyberprzestrzeń i roboty humanoidalne nowym wyzwaniem edukacji*. APS, Warszawa, 2011.
- [2] BRALCZYK J., *Słownik 100 tysięcy potrzebnych słów*. PWN, Warszawa, 2005.
- [3] FRĄCZEK M., MARCZYK M. (red.), *Wybrane aspekty bezpieczeństwa cybernetycznego SZ RP*. AON, Warszawa, 2014.
- [4] JANCZAK J., *Zarządzanie sieciami teleinformatycznymi opartymi na współczesnych technologiach taktycznych WLqđ*. AON, Warszawa, 2011.
- [5] KISIELNICKI J., *Rola informatyki w zarządzaniu*, [w]: BOGDANIENKO J., *Organizacja i zarządzania w zarysie*. WNW UW, Warszawa, 2010.
- [6] KONIECZNY M., *Poszukiwanie tożsamości w cyberprzestrzeni. Implikacje pedagogiczne*. VULCAN, Wrocław, 2012.
- [7] LIDWA W., KRZESZOWSKI W., *Ochrona infrastruktury krytycznej*. AON, Warszawa, 2012.
- [8] MICHNIAK J., *Dowodzenie i łączność*, AON. Warszawa, 2003.
- [9] MICHNIAK J., *Teleinformatyka i technika biurowa*. AON, Warszawa, 2009.
- [10] PACEK B., HOFFMAN R., *Działania Sił Zbrojnych w Cyberprzestrzeni*. AON, Warszawa, 2013.
- [11] WASILEWSKI J., *Zarys definicyjny cyberprzestrzeni*. ABW, Warszawa, 2013.
- [12] WOŁEJSZO J., MARCZYK M., *Identyfikacja i charakterystyka cyberprzestrzeni wykorzystywanej na potrzeby militarne państwa*. AON, Warszawa, 2014.
- [13] WOŁEJSZO J. (red.), *Automatyzacja dowodzenia SZ RP w środowisku sieciocentrycznym*. Centrum Techniki Morskiej, Ośrodek Badawczo-Rozwojowy, Gdynia–Warszawa, 2013.
- [14] WRZOSEK M., *Procesy informacyjne w zarządzaniu organizacją zhierarchizowaną*. AON, Warszawa, 2010.
- [15] Ustawa z dnia 18 lipca 2002 roku o świadczeniu usług drogą elektroniczną, art. 2, ust. 3.
- [16] Ustawa z dnia 15 kwietnia 2005 roku o zmianie ustawy o ochronie informacji niejawnych oraz niektórych innych ustaw.
- [17] Ustawa z dnia 16 lipca 2004 roku Prawo telekomunikacyjne, art. 2. ust. 35.
- [18] Ustawa z dnia 29 sierpnia 1997 roku o ochronie danych osobowych, art. 6, ust. 1.

Cyberspace as a new dimension of human activity – conceptual analysis of the area

ABSTRACT: Cyberspace is a communication space created by Internet connections systems. It allows its users to communicate on the network and establish relationships in real time. Cyberspace is also an information exchange environment using networks and computer systems. In the paper, the author describes the cybernetic space as the dimension of human activity in which all activities diverge in character from the physical environment. It presents the conceptual and structural scope of this space, which, apart from the land, sea, air and space environment, has become a new dimension of human activities, including those of a military nature.

KEYWORDS: cyberspace, ICT network, threats

Praca wpłynęła do redakcji: 16.02.2018 r.

**Informacje dla autorów – procedura recenzowania artykułów
i sposób przygotowania tekstu dla redakcji
PRZEGŁĄDU TELEINFORMATYCZNEGO**

W Przeglądzie Teleinformatycznym zamieszczane są oryginalne artykuły z dziedzin *informatyka, automatyka i robotyka* oraz pokrewnych, np. *telekomunikacja*, niepublikowane dotychczas w całości lub w znaczącej części. Jeśli praca jest fragmentem innej pracy opublikowanej, np. pracy doktorskiej, habilitacji etc., powinna być umieszczona w spisie literatury i redakcja powinna być o tym poinformowana.

W celu opublikowania artykułu w *Przeglądzie* niezbędne jest dostarczenie do redakcji treści artykułu w postaci **elektronicznej** (i ewentualnie papierowej wydrukowanej w jednym egzemplarzu) według podanego szablonu. Przyjmowane są prace w języku polskim lub angielskim. Wydruk powinien być jednostronny, czytelny, na białym papierze formatu A4. Tekst artykułu ma być przygotowany w formacie edytora Microsoft Word (wersje 97–2003, 2007 lub 2010). Szablony dla artykułów są dostępne w pliku na stronie przeglad.ita.wat.edu.pl (lub review.ita.wat.edu.pl). Przekazane materiały do redakcji powinny zawierać plik tekstu w formacie *.doc (lub *.docx) ze wstawionymi rysunkami. Redakcja nie przepisuje tekstów i nie wykonuje rysunków. Oprócz pliku *.doc (lub *.docx) należy dostarczyć pliki źródłowe rysunków (najlepiej w formacie *.eps, *.jpg, *.tif lub innym powszechnie używanym standardzie).

Artykuły przeznaczone do opublikowania w *Przeglądzie* podlegają wstępnej ocenie przez redaktora działu, a następnie podlegają recenzji przez dwóch zewnętrznych recenzentów. Recenzenci i autorzy nie znają swoich danych personalnych. Z treścią recenzji można zapoznać się w redakcji. Jeśli jedna z recenzji jest negatywna (lub nieprecyzyjna), może być powołany trzeci recenzent. Jeśli dwie recenzje są negatywne, artykuł nie jest publikowany. Jeśli z recenzji wynika konieczność dokonania poprawek w treści artykułu, to wykonuje je autor i dostarcza do redakcji poprawioną wersję artykułu w terminie ustalonym przez redakcję.

Objętość artykułu zasadniczo nie powinna przekroczyć 20 stron maszynopisu A4. Odstąpienie od tej zasady wymaga uzgodnień z redakcją *Przeglądu*. Na stronie tekstu artykułu nie może być pozostawione więcej niż 10% pustego miejsca. Rysunki powinny być numerowane (podpis pod rysunkiem), tabele również numerowane (tytuł na górze tabeli). Literatura może być uszeregowana alfabetycznie i powinna mieć postać jak w szablonie.

Autor (autorzy) składają do redakcji oświadczenia autorskie o wkładzie procentowym w powstanie artykułu, braku wcześniejszej publikacji artykułu w przedstawionej formie lub wystąpieniu publicznym na ten temat na konferencji lub sympozjum.

Redakcja zastrzega sobie prawo wprowadzenia niewielkich redakcyjnych zmian w treści artykułu bez konsultacji z autorem. Redakcja prosi, aby **nie stosować** żadnego specjalnego formatowania i **trzymać się ściśle** ustaleń zawartych w szablonie.

Streszczenia i pełne teksty artykułów drukowanych w *Przeglądzie* zamieszczane są w krajowej bazie danych o zawartości polskich czasopism technicznych BazTech oraz na platformie INDEX COPERNICUS. Opublikowane w *Przeglądzie* artykuły będą także w całości udostępnione w internetowej wersji czasopisma pod adresem przeglad.ita.wat.edu.pl (lub review.ita.wat.edu.pl).

Publikacja w *Przeglądzie* nie wiąże się z żadnymi kosztami dla autorów. Redakcja nie pobiera opłat za zgłoszenie, przygotowanie do druku i publikację pracy. Przekazanie artykułu do publikacji w *Przeglądzie* jest równoznaczne z przekazaniem autorskich praw majątkowych do publikacji na rzecz wydawcy – Wojskowej Akademii Technicznej.



Wszystkie artykuły opublikowane w czasopiśmie **PRZEGŁĄD TELEINFORMATYCZNY (TELEINFORMATICS REVIEW)** są udostępniane na licencji Creative Commons Uznanie autorstwa – Użycie niekomercyjne – Bez utworów zależnych 3.0 (CC BY-NC-ND 3.0), która zezwala na kopiowanie, przedstawianie i rozpowszechnianie utworu jedynie w celach niekomercyjnych oraz pod warunkiem zachowania go w oryginalnej postaci (czyli nietworzenia utworów zależnych), przy jednoczesnym odpowiednim oznaczeniu autorstwa utworu.

Redakcja nie zwraca materiałów dostarczonych do redakcji.

Redakcja nie przewiduje honorariów za opublikowanie artykułu.

Redaktor naczelny może odmówić opublikowania materiałów autorskich w przypadku, gdy:

- treści zawarte w materiałach naruszają prawo (zasady ochrony tajemnicy, prawo prasowe, prawo autorskie itp.) lub dobre obyczaje;
- autor nie zgadza się na wprowadzenie wszystkich koniecznych poprawek zaproponowanych przez redakcję lub recenzentów;
- tekst i materiał ilustracyjny złożony przez autora nie spełnia wymagań technicznych podanych w niniejszym dokumencie i szablonie.