

Advanced Phishing Techniques: Analyzing Adversary-in-the-Middle and Browser-in-the-Browser Attacks in Modern Cybersecurity

Eric B. Blancaflor, Jacob O. Duldulao, John Vincent E. Espeño, Geoff Stephen M. Patag, Ma. Theresa Menor, Grace Lorraine Intal

School of Information Technology, Mapúa University, Makati, Philippines

E-mails: ebblancaflor@mapua.edu.ph joduldulao@mymail.mapua.edu.ph
jveespeno@mymail.mapua.edu.ph gsmptag@mymail.mapua.edu.ph mtmarquez@mapua.edu.ph
glintal@mapua.edu.ph

Abstract: Phishing attacks remain among the most common techniques malicious actors use to steal sensitive information. This paper examines two emerging phishing techniques: Adversary-in-The-Middle (AiTM) and Browser-in-the-Browser (BitB) attacks. AiTM attacks intercept communicating devices, allowing attackers to hijack accounts and access sensitive data. BitB attacks use a deceptive pop-up login window that mimics a legitimate authentication portal, forcing users to input private credentials. These methods have been developed to bypass traditional security measures, especially Multi-Factor Authentication (MFA), posing an ever-growing threat to real-life sectors such as finance, healthcare, and public services. These attacks are becoming more prevalent across various sectors, calling for businesses to implement stronger security measures. Effective countermeasures include detection and prevention, mitigation to limit attack impact, and AI-based attack identification and termination tools. Organizations can reduce the risk of these sophisticated cyber threats through a combination of prevention, mitigation, and AI-based tools.

Keywords: Adversary-in-The-Middle (AiTM), Browser-in-the-Browser (BitB), Phishing, Multi-Factor Authentication (MFA), Social engineering, Zero Trust security, Ai-Driven threat detection.

1. Introduction

1.1. Context of the study

Phishing attacks continue to be among the most common techniques utilized by malicious actors to acquire sensitive data such as passwords or credit card details. According to the 2023 Verizon Data Breach Investigations Report (DBIR) [1], phishing remains a leading cause of security breaches, contributing significantly to data theft incidents across various industries. Despite advancements in security

technologies, including the widespread adoption of Multi-Factor Authentication (MFA), threat actors have adapted by evolving their phishing methods to overcome these defenses. Reports from ThreatX and other cybersecurity firms [2] highlight that modern phishing attacks now employ sophisticated tactics, enabling threat actors to bypass traditional security measures like MFA. This paper focuses on two such advanced techniques: Adversary-in-The-Middle (AiTM) and Browser-in-the-Browser (BitB) attacks, which illustrate how phishing strategies continue to evolve in response to enhanced security protocols.

Historically, phishing attacks primarily relied on deceptive emails and fraudulent websites to acquire credentials. However, contemporary attack methods such as AiTM and BitB enable threat actors to circumvent even advanced protections like Multi-Factor Authentication (MFA). For instance, an AiTM attack positions the malicious actor between the user and a legitimate website to intercept authentication credentials and session tokens, enabling unauthorized session control even with MFA implemented. Research from Desclope demonstrates how attackers can capture these tokens to maintain persistent access to compromised accounts presenting significant security risks [3]. Similarly, BitB attacks employ sophisticated interface replication techniques within legitimate websites, creating authentic-appearing login interfaces that are challenging to distinguish from legitimate authentication windows [4].

This paper examines the operational mechanisms and effectiveness of AiTM and BitB attacks, highlighting examples from sectors such as finance and healthcare, which are common targets due to the high value of their data. A 2023 analysis by Microsoft indicates how AiTM techniques enable threat actors to stage additional attacks, such as business email compromise, potentially disrupting financial transactions and data handling in these sectors [5]. Such incidents illustrate the broad impact these attacks have across different industries, resulting in both data breaches and financial losses.

To counter these evolving phishing techniques, it is critical to explore defenses that can minimize risk. Solutions such as zero-trust security models, advanced AI-driven threat detection, and enhanced user awareness training may enable organizations to maintain defensive superiority. Zero Trust architecture, for example, implements continuous verification of user and device identity throughout each session, rather than relying solely on initial authentication [6]. Given the rapid evolution of phishing techniques, implementing proactive security measures is essential for protecting users in the contemporary digital environment.

1.2. Background on phishing attacks

Phishing represents a form of social engineering attack where malicious actors attempt to acquire sensitive information by impersonating legitimate entities [7]. The term “phishing” comes from the analogy of using bait to catch targets, reflecting how attackers use deceptive lures to capture credentials and sensitive data [8]. According to cybersecurity research, the first phishing attacks started in the 1990s, primarily targeting financial institutions through email-based deception [9].

As defensive technologies evolved, attack methodologies adapted accordingly. Instead of just sending emails, threat actors developed fake websites, sent harmful

files, and even called people on the phone to deceive them [10]. Research indicates that phishing still works well today because it exploits human trust factors, which makes it hard even for technically proficient people to spot fake messages [1].

When security tools got better at stopping basic phishing emails, attackers changed their methods. The threat actors developed spear phishing techniques, where they target specific people or companies using information they find online about them [11]. Recent studies show that modern phishing is not just about stealing passwords anymore - malicious actors now use it to deploy malware, steal money, or break into important computer systems [12].

Even though we have better security tools now, like email filters and firewalls, phishing is still one of the most successful ways to attack people and companies. This is because it exploits human trust, which technical solutions cannot fully protect against [13]. When companies started using Multi-Factor Authentication (MFA) to stop phishing, attackers came up with new methods like Adversary-in-the-Middle (AiTM) and Browser-in-the-Browser (BitB) attacks to get around these protections [14].

Research shows that these new phishing attacks are getting harder to spot and more focused on specific targets [15]. This means both regular people and companies need to keep learning about new ways attackers might try to deceive them.

1.3. Objectives of the study

This research aims to achieve three main objectives:

1. To examine and analyze two emerging phishing attack methods: Adversary-in-the-Middle (AiTM) and Browser-in-the-Browser (BitB). This study will investigate how these attacks circumvent traditional security measures, particularly Multi-Factor Authentication (MFA), and document their growing sophistication.
2. To investigate the impact of AiTM and BitB attacks on critical sectors, specifically banking, healthcare, and government institutions. Through case study analysis, this research will demonstrate how these attacks are executed and their consequences, including financial losses, data breaches, and service disruptions.
3. To evaluate and present effective defense strategies against these advanced phishing techniques. This includes exploring modern security approaches like Zero Trust architecture and AI-based threat detection, as well as assessing the role of employee training in preventing successful attacks.

Through these objectives, this study aims to provide a comprehensive understanding of modern phishing threats and their countermeasures. By analyzing both the technical aspects of these attacks and their real-world impacts, while also examining defensive strategies, this research will contribute to the broader field of cybersecurity by helping organizations better protect themselves against evolving phishing techniques.

1.4. Scope and methodology of the study

This study employs a descriptive approach to examine AiTM and BitB attacks, integrating academic literature analysis with real-world security reports to analyze these attack methodologies and their organizational impact.

The researchers gathered information from academic databases such as IEEE and ACM, along with technical reports from leading technology organizations such as Microsoft, Google, and IBM. The study also incorporated technical guidelines from the National Institute of Standards and Technology (NIST) and government security advisories. Utilizing the snowball sampling methodology, the researchers followed reference lists from key papers to identify additional relevant sources.

The literature review focused on materials from 2020 to 2024, prioritizing recent developments in phishing attacks. Source selection emphasized work from established researchers and organizations, combining both technical analysis and documented attack cases.

The research examines four main areas:

- Technical mechanisms: How AiTM and BitB attacks operate.
- Attack impact: Effects across different sectors.
- Defense strategies: Current countermeasures and effectiveness.
- Future outlook: Emerging threats and trends.

This comprehensive approach provides clear insights into modern phishing techniques and their implications.

2. Literature review

2.1. The role of Multi-Factor Authentication (MFA) in security

MFA is a security mechanism that adds additional layers of protection when logging into accounts. Instead of just using a password, users need to prove their identity in two or more ways [16]. These methods typically include something the user knows (such as a password), something they possess (such as a mobile device), or something inherent to them (such as a fingerprint).

Companies started using MFA because passwords alone are not safe enough. Many people use the same password for different accounts or choose passwords that are easy to guess. Even if credentials are compromised through phishing or data leaks, unauthorized access remains blocked without the secondary authentication factor, such as a code sent to a mobile device [10].

MFA makes accounts much safer, which is why many companies use it. Big companies like banks and email services often make their users turn on MFA to keep their accounts safe from hackers. The extra security step increases the difficulty for malicious actors to break into accounts even with compromised passwords.

However, some attackers have found ways to get past MFA. They use sophisticated techniques like AiTM attacks [17]. These attacks work by catching both the password and the security code at the same time when someone tries to log in. This enables unauthorized access to accounts even when MFA is turned on, which is a big problem for security experts.

According to the description of MITRE, adversaries manipulate authentication mechanisms by altering authentication processes, such as those managed by LSASS and SAM on Windows, PAM on Unix-based systems, and authorization plugins on macOS. The modification of these processes can help attackers bypass security controls, extract credentials, or maintain persistent access to remote systems and

services like VPNs, Outlook Web Access, and remote desktops. This technique allows them to move laterally within a network and exploit sensitive resources without triggering standard authentication safeguards [18].

As more people use MFA, attackers keep developing sophisticated techniques like AiTM and Browser-in-the-Browser (BitB) attacks to fool users [18]. This shows that while MFA helps with security, it cannot protect against everything by itself. Companies need to use MFA along with other security tools to better protect their users from these new attacks.

Table 1. MFA strengths and weaknesses

MFA method	Strengths	Weaknesses
SMS-based authentication	Easy to set up and use for most users	Vulnerable to SIM swapping and phishing attacks
		Relies on mobile network availability
Mobile app (authenticator)	More secure than SMS (codes generated on the device)	Vulnerable to AiTM attacks (real-time interception)
	Does not rely on network service	May require internet access for the initial setup
Email-based authentication	Familiar to users	Email accounts can be compromised, leading to account recovery or access issues.
	Easy to implement	Vulnerable to phishing attacks
Hardware tokens	Strong physical security	Costly to implement and distribute
	Difficult for attackers to replicate	Risk of being lost or stolen
		Inconvenient for users without a device
Biometric authentication	Adjusts authentication requirements based on user behavior	May still rely on passwords and be vulnerable to advanced phishing or AiTM attacks
	Less intrusive for users	
Push notifications	Convenient for users	Can be exploited in AiTM attacks by intercepting approval
	Requires approval via mobile device	Depending on the availability of the device
Risk-based authentication	Adjusts authentication requirements based on user behavior	May still rely on passwords and be vulnerable to advanced phishing or AiTM attacks
	Less intrusive for users	
Single Sign-On (SSO)	Simplifies authentication across platforms	If the SSO provider is compromised, all linked accounts are at risk
	Can integrate with MFA	Adoption can be limited without widespread integration

2.2. Analysis of adversary in the middle attacks

2.2.1. Adversary in the middle attacks

The man-in-the-middle attack is a common type of network attack wherein an intruder can sit in between two communicating network endpoints, gaining access to

data being transferred throughout the communication. MitM attacks come in forms such as [19]:

- Passive MitM
 - Eavesdropping on a conversation to access information being sent throughout.
- Tampering
 - Editing information being shared between two endpoints.
- Delaying
 - Delaying data being sent between two networks, causing an interruption in the network's processes.
- Dropping
 - Completely deleting information being sent during communication, causing significant data loss within the network.

Various approaches are also used to penetrate a connection and infiltrate a communication, such as spoofing and decryption. Anyone can become a victim of man-in-the-middle attacks, as communication between multiple devices is common in several types of system architectures and applications. Upon breach of communication, attackers can access information and remain almost unnoticed, making MitM prevention a serious security concern for an organization.

A widespread type of man-in-the-middle attack is the adversary-in-the-middle attack. Adversary-in-the-middle attacks can refer to the active versions of man-in-the-middle attacks, wherein information is hijacked and stolen, or an attacker interacts with information while it is being sent throughout the network. The adversary in the middle attacks is a term normally related to multi-factor authentication bypassing, as more prevalent approaches have found ways to infiltrate these authentication processes.

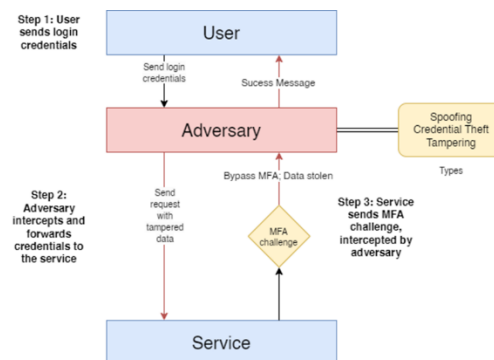


Fig. 1. Diagram representation of an AiTM attack

2.2.1. AiTM attacks against MFA

Adversary in-the-middle attacks have been rampant against MFA. This is because MFA follows a decentralized, multi-device approach, which requires the user to log in using more than one device to approve verification. As this may be a better approach to improve simpler authentication such as password-based authentication in preventing various attacks such as password hijacking and shoulder surfing, its

decentralized nature makes it vulnerable to AiTM. Based on the AiTM model, attackers can intercept communication between two devices, allowing them to illegally phish for confidential information and real-time authentication data, without being detected. With this, multi-factor authorization must be properly studied and implemented to ensure its overall security.

A study by Amft et al. [20], aimed to assess account recovery loss using multi-factor authentication in terms of its security, implementation, and user experience. With the rapid growth of MFA as a main form of authentication to improve the simpler authentication methods such as password-based authentication, its drawbacks must also be assessed and require equal attention. In the study, 71 websites were assessed in their use of MFA to determine how insecure the use of MFA is, along with determining steps to take to improve its implementation. The researchers conducted the study by creating accounts on websites that required MFA and utilizing the MFA to recover the account after a certain amount of time. Through this, the researchers were able to determine that mobile apps, SMS, emails, and hardware tokens were the most used methods of MFA. Not only this but some websites have also been assessed as implemented poorly, as they were unable to gain access to 23 of the accounts.

As multi-factor authentication was developed to strengthen password-based authentication, according to Gavazzi et al. [21], MFA has been recorded with low adoption rates. Because of this, risk-based authentication has become another highly recommended form of authentication. This type of authentication assesses the probability of account compromise, and if detected, prompts the user for more verifying action to gain access. The study by Gavazzi, et. al. aimed to measure the availability and usage of MFA and RBA on the web, along with additional authentication factors used, and the use of single sign-on across various websites. 208 popular sites in the Tranco top 5K that support account creation were used to assess the study. Based on the study, only 43% of the websites audited offered any form of MFA. Upon further assessment, however, if each account that does not support MFA and/or RBA were to be made through an SSO provider that does, about 80% would have access to MFA and 72% of sites would block suspicious login attempts.

2.2.3. Ettercap and Wireshark on performing adversary in the Middle attacks

Ettercap and Wireshark are widely used tools used in simulating and performing man-in-the-middle attacks. Ettercap, mainly compatible with Unix operating systems, is a comprehensive suite that compiles many features that can be used for host analysis, while Wireshark is a tool that can be used to analyze network traffic in detail. Ettercap is a program in the Kali Linux operating system that performs Address Resolution Protocol poisoning [22]. This is a free and open-source network security tool for Man-in-the-Middle (MitM) attacks, protocol analysis, and network traffic manipulation. According to the Ettercap official website, it can perform acts such as protocol analysis, traffic manipulation, SSL/TLS decryption, password sniffing, and DNS spoofing. Wireshark, on the other hand, is a flexible network protocol analyzer that plays a crucial role in capturing and analyzing network traffic and supporting network administrators in resolving network-related problems [23]. Recognizing

itself as the world's foremost network protocol analyzer, the tool features various functions such as deep inspection, offline analysis, rich VoIP analysis, decryption support, and coloring rules.

According to Cekerevac et al. [24], MitM can be initiated through the following approaches:

- ARP cache poisoning – attackers intercept ARP tables to redirect network traffic between two parties.
- DNS spoofing – attackers forge DNS responses to intercept information.
- Session hijacking – an act of impersonation using the user's token to gain access and perform actions.
- SSL hijacking – attackers utilize fake certificates to decrypt SSL-encrypted communication.

Through Ettercap, researchers can simulate ARP cache poisoning. Ettercap enables users to manipulate and transmit ARP packets, while Wireshark allows monitoring of these packets throughout the node communication. This combination enables researchers to implement and simulate an adversary-in-the-middle attack, along with being able to record and manipulate this process through Wireshark. Modern studies have turned to using Ettercap and Wireshark together to conduct man-in-the-middle simulations. By using Ettercap, users can perform traffic interception and manipulation, while Wireshark can be used for protocol analysis and packet inspection.

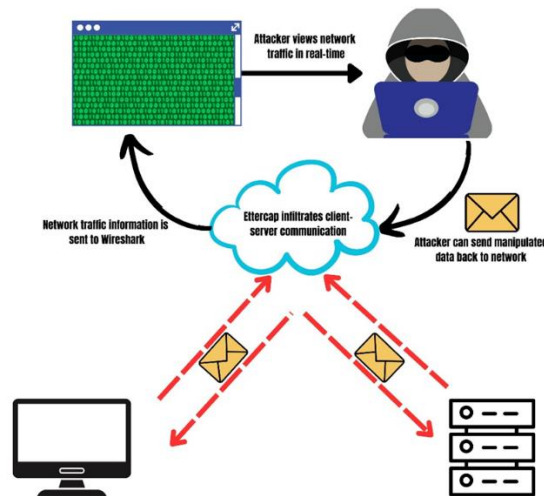


Fig. 2. Ettercap and Wireshark in an AiTM attack

According to Rajendran [22], the SSL is prone to MitM attacks since the process of establishing communications is unaware of who is being connected. Through the SSL three-way handshake, the connection does not explicitly determine who is connecting with the server or the client, making it vulnerable to attackers intercepting the communication. Though Certificate Authority (CA) can be used to thwart MitM attacks within the SSL protocol, hijacking methods are still able to

create fake certificates to penetrate the connection. The study by Rajendran [22], aims to determine ways to improve MitM detection during secure encryption communication. To do this, the researchers measured the time it takes to complete an SSL handshake, along with looking at the typical traffic patterns throughout the process. Kali Linux, Ettercap, and MitM proxy tools were set up to compare data collected from the Wireshark records captured on the client machine.

Through Ettercap, the researchers were able to simulate ARP poisoning conducted by a Linux device setup in VM Oracle Box to penetrate a Windows 10 virtual machine. Wireshark is then installed on the client machine to capture the packet of the TCP and SSL handshake between the client and the web application. Using Wireshark filter options, the study captured the TCP and SSL steak flow with the sequence timestamps for each website. With this setup, the response time of 30 websites across the globe was analyzed, and the average RTT for TCP and SSL were calculated. Based on the results, it was shown that all websites take a longer response time after a MitM attack. The consistently perceptible timing differences for SSL/TLS connections from the 30 samples collected can be used as key parameters to detect the existence of MitM threats.

A study by Chavoshi et al. [23], aimed to simulate a MitM attack on a cyber-physical system to evaluate the effects of the attack on control systems against MitM threats. According to the study, ARP is a pivotal enabler for MitM attacks, as ARP enables the discovery of the MAC address of a node. The act of ARP poisoning aims to compromise the ARP table of a target by associating the IP address of the counterpart with the MAC address of the attacker. This allows the attacker to execute the MitM attack. With this, an attack could be detected based on an irregular behavior presented by the system, which can be detected by the volume and nature of the transmitted and received information. The primary purpose of this article is to explore and implement MitM attacks, alongside using machine learning classification algorithms to detect these attacks and their potential.

Table 2. Common technologies in an AiTM Attack implementation

Technology	Function	Application in AiTM	Unique features	Limitations
Ettercap	Network security tool focused on MitM attacks, protocol analysis, and traffic manipulation	ARP poisoning, DNS spoofing, SSL hijacking, and session hijacking	Open-source, ARP packet manipulation, support for SSL/TLS decryption, protocol-specific attacks	Primarily for Unix OS, limited advanced packet analysis without Wireshark
Wireshark	Network protocol analyzer used for packet inspection and protocol analysis	Captures and monitors network traffic during MitM attacks, e.g., inspecting TCP and SSL handshakes	Deep packet inspection, VoIP analysis, SSL/TLS decryption support, custom filters for protocol inspection	Passive monitoring only requires integration with attack tools like Ettercap for traffic manipulation

The study was conducted under a cyber-physical liquid-level control system. Utilizing a Networked Control System (NCS) with a PI controller, the system communicated via Wi-Fi and TCP/IP protocol. With this setup, remote monitoring and control were enabled, creating an ideal environment for AiTM to be simulated, as the devices' communication can be controlled from a foreign source caused by an interception [3]. Ettercap was mainly utilized for carrying out MitM attacks through ARP poisoning, while Wireshark was used to monitor the packets transmitted from the computer to the datalogger to disrupt the control signal and output of the sensor [23].

The primary purpose of the mentioned studies was to explore and implement MitM attacks, alongside using machine learning classification algorithms to detect these attacks and their potential. The study by [22] simulated the performance of AiTM on the SSL handshake, being able to intercept the communication between a client and server. The study by [23] conducted AiTM on a liquid-level network-controlled system that enabled devices to communicate through the Internet. Both studies followed a similar setup of a decentralized architecture, where devices were separated and relied on a network to communicate. Through this, an AiTM attack can be conducted on this type of structure, which calls for an improvement in security for decentralized networks and architectures.

2.2.4. Mitigation and controls

Understanding how to protect against Adversary-in-the-Middle attacks requires a multi-layered approach to security. Research by G a v a z z i et al. [21] indicates that while 43% of websites offer MFA, implementing additional security layers significantly improves protection against sophisticated attacks.

Network security controls – modern network security must address both traditional and emerging threats. Key controls include:

- Network segmentation using VLANs and microsegmentation;
- Implementation of TLS 1.3 with Perfect Forward Secrecy (PFS);
- Real-time Deep Packet Inspection (DPI) for traffic analysis;
- Next-generation firewalls with application-level filtering;
- Integration of Security Information and Event Management (SIEM) systems.

According to A m f t et al. [20], certificate-based authentication shows particular promise in preventing AiTM attacks. Their study of 71 websites revealed that hardware tokens and certificate-based methods provided the strongest protection against session hijacking attempts.

Authentication and session management – research by C e k e r e v a c et al. [24] emphasizes the importance of robust session management. Critical components include:

- X.509 client certificates for mutual authentication;
- Session token rotation every 15-30 min;
- Implementation of OAuth 2.0 with PKCE (Proof Key for Code Exchange);
- Secure token storage using HTTP-only cookies with SameSite=Strict;
- Real-time monitoring of session characteristics using machine learning models.

Monitoring and response – building on research by Kusumo, Erlangga and Ramadhan [32], effective monitoring should include:

1. Network analysis
 - a. Machine learning-based traffic pattern analysis;
 - b. Behavioral analytics for user session profiling;
 - c. Real-time SSL/TLS certificate validation;
 - d. Automated response to detected anomalies.
2. Authentication monitoring
 - a. Geographic location analysis for login attempts;
 - b. Device fingerprinting and reputation checking;
 - c. Time-based access pattern analysis;
 - d. Multi-point session validation.

User awareness – studies by Denbigh-White and Ventura [2] emphasize the critical role of user education. Training programs should include:

- Practical phishing simulation exercises;
- Certificate validation workshops;
- Security incident reporting procedures;
- Regular security awareness updates.

Implementation strategy – organizations should follow a phased approach to control implementation:

1. **Phase 1.** Foundational security
 - a. Basic network segmentation;
 - b. TLS 1.3 deployment;
 - c. Initial monitoring setup.
2. **Phase 2.** Enhanced control
 - a. Advanced authentication methods;
 - b. Behavioral analytics;
 - c. Automated response systems.
3. **Phase 3.** Advanced protection
 - a. Zero Trust architecture;
 - b. AI-driven threat detection;
 - c. Continuous security validation.

This comprehensive approach aligns with findings from Microsoft's Digital Defense Report [10], which highlights the importance of layered security measures in combating modern AiTM attacks. Regular assessment and updating of these controls ensure continued effectiveness against evolving threats.

2.3. Browser-in-the-Browser (BitB) attacks

2.3.1. What are BitB attacks

BitB attacks represent a new type of attack vector where attackers make fake login windows that look exactly like real ones within web browsers [25]. These fraudulent windows precisely mimic authentic login interfaces from trusted providers such as Google or Microsoft. The deceptive interfaces are designed to harvest user authentication credentials.

The typical attack flow occurs when users attempt to utilize “Sign in with Google” or similar single sign-on services. The attack presents users with a fake authentication window that replicates legitimate login interfaces with high fidelity [26]. Most people do not notice anything wrong because they see these login windows all the time. When they type in their username, password, and security code, the malicious actors behind the attack can steal all this information and use it to break into their accounts.

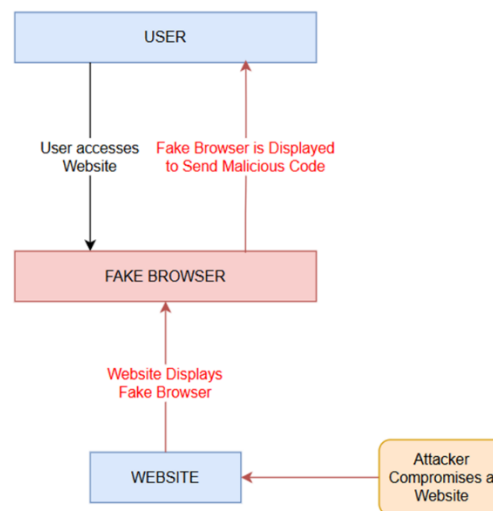


Fig. 3. BitB attack model

2.3.2. Technical aspects of BitB attacks

BitB attacks work well because attackers can make fake login windows that look exactly like real ones [27]. These fake windows copy everything from the real login screens – the way they look, the text style, and even the colors of buttons. Attackers are very careful to make everything look real, including adding a fake website address at the top (like **accounts.google.com**) and security symbols like padlocks. This makes it very hard for people to spot that something is wrong.

From a technical perspective, these fake windows are not actually separate browser windows. Instead, they’re built directly into the webpage using special computer code (HTML and JavaScript). While real pop-up windows can be moved around or made bigger, these fake ones cannot – though most people do not notice this difference [28]. This method helps attackers avoid security tools that usually block suspicious pop-up windows.

BitB attacks often target something called Single Sign-On (SSO), which lets people use one account (like their Google account) to log into many different websites [16]. Because people are used to seeing these login windows appear, they typically proceed without additional verification. This makes the attack more likely to work.

When someone types their username and password into the fake window, the attacker can steal this information right away. If the person also enters a security code

(like from their phone), the attacker can grab that too and use it to break into the account immediately.

2.3.3. Examples of BitB attacks in real-life applications

BitB attacks have been found targeting many different services, especially ones that let you log in with accounts like Microsoft 365, Google Workspace, or social media [10]. These attacks work well because people are used to seeing login windows pop up when they use these services.

A big example of these attacks happened to Microsoft 365 users. When employees clicked on links in fake emails, they saw what looked like a normal Microsoft login window. The window looked so real that people typed in their usernames, passwords, and security codes without knowing they were giving this information to attackers. Once attackers got in, they could read company emails and steal important files [29]. According to security reports, this kind of attack led to several big data breaches in 2023.

Another serious case happened at a financial institution where attackers used BitB techniques to steal login details from bank employees [30]. They made a fake “Sign in with Google” window that looked legitimate. When bank workers tried to log in through this fake window, attackers stole their Google account information. This let the attackers get into the bank accounts and steal money.

These real examples show how dangerous BitB attacks can be for both regular people and big companies. They’re especially bad for places that deal with important things like money, health information, or business secrets, where one successful attack can cause significant operational impact.

2.3.4. Browser Exploitation Framework (BeEF) and Cross Site Scripting (XSS) in BitB attacks

A study by Kusumo, Erlangga and Ramadhan [32], aimed to investigate the types of vulnerabilities and attacks that affect Infrastructure as a Service (IaaS) and Platform as a Service (PaaS) cloud models. The methodology involved setting up these attacks in a controlled environment to observe how they exploit cloud-related vulnerabilities. The MitB attack used XSS Hooking and JavaScript Injection methods to manipulate the user’s browser. The attackers injected malicious scripts that enabled them to steal sensitive information such as usernames and passwords. Through their study [32], they were able to display a fake Google Mail login page to steal the victim’s credentials, where these credentials were captured and displayed in the BeEF interface.

A MitB attack is conducted by utilizing an agent of security breach that can alter the connection between two parties through the means of the target’s browser [32]. Similarly, a BitB attack breaches a target’s machine through the browser to deceive users and mainly focuses on deception and social engineering [27]. The attacker simulates a legitimate browser window to mimic login prompts, deceiving users into providing sensitive information. That said, studies that implement and simulate MitB attacks, such as this, can also be used to raise awareness and study possible mitigation strategies against BitB attacks.

Table 3. Descriptive comparison between BeEF and XSS

Browser Exploitation Framework (BeEF)	Cross Site Scripting (XSS)
A penetration testing tool that focuses on the web browser	A type of attack used to perform MitB attacks
Allows a profession to assess the security posture of a target environment by using client-side attack vectors	Injects malicious code to legitimate websites that contain malicious code and scripts
Examines exploitability within the content of the web browser	Injects malicious code throughout sections of a website and the most common network attacks via the web

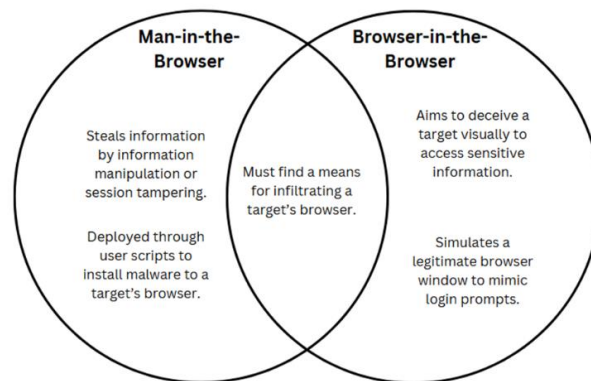


Fig. 4. Comparison of MitB and BitB attacks

A study by Alsaffar et al. [31], aimed to develop a software program capable of detecting various types of XSS attacks. This study aimed to identify possible threats brought by this technology, along with recommendations on how to handle these threats. BeEF and XSS can be used together to perform various attacks such as MITB and BITB, and looking into the mitigation of threats from XSS can be an effective approach in reducing such attacks. The study introduces a software program developed using the Delphi programming language, designed to map web applications entirely and detect XSS vulnerabilities in both public and private (authenticated) sections of websites. The software was tested against web applications and demonstrated a 99.47% detection accuracy with high precision and recall rates. Compared to existing tools, it spent 44% less time finding vulnerabilities than Acunetix and 20% less time than XSpider.

The studies mentioned aimed to analyze the implementation of BeEF and XSS to conduct MitB attacks. Points where malicious scripts can be injected were analyzed, along with discovered vulnerabilities. In the study, they conducted an MitB attack using both BeEF and XSS to gain credential information through a fake Google Mail login page.

A fraudulent browser can be carried into any part of a website, especially authenticated ones, so these applications can be studied as a way of mitigating a BitB attack. Studies that implement and simulate MitB attacks, such as this, can also be used to raise awareness and study the possible mitigation strategies against BitB attacks. Improving the detection of MitB vulnerabilities can strengthen defenses against both MitB and BitB attack strategies.

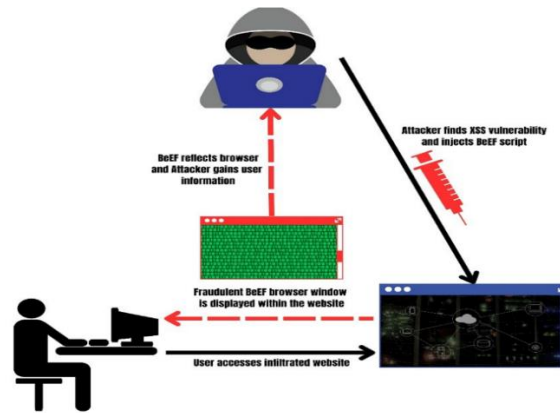


Fig. 5. XSS and BeEF BitB Attack model

Table 4. Technologies enabling BitB attacks

Technology	Function	Application in BitB	Features	Limitations
BeEF	Browser exploitation and penetration testing	Simulate legitimate browser windows to deceive users	High customization, integration with various scripts	Limited to browser-based vulnerabilities
XSS	Injection of malicious scripts into web pages	Can insert fraudulent interfaces within a trusted site	Allows control over user experience within the browser	Requires vulnerabilities in the target site

3. Case studies

3.1. Case 1. An AiTM attack on the financial sector

A serious example of an Adversary-in-the-Middle (AiTM) attack was conducted at some well-known banks [33]. The attack progression was as follows:

- **Initial Vector.** Attackers sent targeted phishing emails to bank employees.
- **Attack Method.** Used sophisticated proxy systems to intercept login credentials and session tokens.
- **Impact.** Unauthorized access to customer accounts and financial transactions.
- **Losses.** Significant financial losses through fraudulent transfers.
- **Detection.** The Attack was discovered only after customers reported suspicious activities.
- **Key Vulnerability.** Reliance on traditional MFA without additional session monitoring.

The attackers sent fake emails to bank workers and customers, tempting them to click on fraudulent links. These links led people to fake bank websites that looked exactly like the legitimate ones. People typed in their usernames, passwords, and security codes without knowing they were on a fake site.

When someone entered their login details, the attackers did something clever: they covertly passed this information to the real bank website right away [34]. This allowed them to acquire special access codes (called session tokens) from the real site. With these codes, they were able to access people's bank accounts even when these accounts were protected by advanced security measures. A critical security concern was that all activities appeared normal to users accessing their accounts.

This type of incident is usually classified as Real Estate Business Compromise (RE-BEC), an infamous type of AiTM attack comprising of impersonation, malware, and spear phishing tactics to launder money from larger financial institutions [35]. According to security reports, one bank experienced significant financial losses when attackers used this technique to get into customer accounts and transfer money around [35]. The bank only found out about the attack after customers started reporting strange activities with their accounts. By then, the laundered money had been fully transferred through innumerable unique accounts making it very difficult to retrieve.

Other instances also involved the use of a money mule, or an individual sent by attackers to initiate or handle unpermitted money transfers to conduct a successful launder. These mules, at times, can be unbeknownst to the situation, such as when attackers use social engineering to manipulate an individual, alongside AiTM strategies. Money mules can be used as a stepping stone for an attacker to conduct fatal phishing attacks. Some instances involve the use of a money mule to create a new bank account, serving as an access point for attackers to compromise a financial institution [35]. With this, the facilitation of a RE-BEC attack is simplified by gaining access to relevant email accounts, hiding an attacker's identity, and legitimizing fraudulent transactions, further complicating the retrieval processes institutions will have to perform. Thus, the combination of social engineering, spear phishing, impersonation, and malware makes AiTM attacks a powerful and grave threat to any financial institution.

This real case shows how dangerous AiTM attacks can be, especially since they can get past security measures that usually work well. It also shows why banks need to implement more secure systems to identify unusual account activity, even when someone seems to be logging in ordinarily.

3.2. Case 2. A BitB attack on the healthcare sector

A serious BitB attack happened recently in healthcare, where attackers targeted doctors and nurses who used online systems to manage patient information [36]:

- **Initial Vector.** Fake emails impersonating hospital software provider.
- **Attack Method.** Sophisticated BitB attack mimicking Google SSO login.
- **Impact.** Compromise of patient medical records and personal information.
- **Consequences.** HIPAA compliance violations and disruption of medical services.
- **Detection.** Discovered through routine security audit.
- **Key Vulnerability.** Insufficient user training on identifying fake login windows.

The attackers sent fake emails that seemed as if they were sent from the hospital's software company. When medical staff accessed these links from these emails, they saw what looked like a normal login page with a Sign-in with Google window. This Google login window was actually fake, made using BitB attack methods. It looked exactly like a real Google login screen, with all the usual security signs that medical staff were used to seeing. When the healthcare workers put in their Google usernames, passwords, and security codes, they were unaware that they were actually giving this information to attackers [29].

Hospitals usually use cloud-based systems for electronic health records and communications, and compromising a single account often provides attackers access to the much broader system. In this instance, the attackers moved within the system with the use of the stolen credentials, granting them entry to administrative privileges. This allowed them to access and modify patient records, along with injecting even more malware into critical software. This type of information is usually sought after by attackers to be stolen as either ransom or sold on dark web marketplaces.

Once the attackers got into these Google accounts, they were able to steal private patient information, including medical records and personal details. This caused big problems for the hospital – forcing them to hinder operations while they figured out what happened. The delay in detecting the breach implied that the sensitive patient data had already been tampered with or encrypted. Because they lost control of patient information, the hospital also got in trouble for breaking HIPAA rules, in the context of protecting patient privacy [37].

Further, the compromised accounts were used to forward further phishing emails to other hospital staff and partner organizations, making the breach spread even wider. Since the phishing emails were sent from official accounts, employees and external staff had no reason to deny its' legitimacy, further increasing the attacker's reach. The damage from this instance resulted in significant financial loss for a whole web of medical institutions and organizations [37].

This attack shows how dangerous BitB attacks can be, especially in places like hospitals that store sensitive information. It also shows why it is so important to teach healthcare workers about these attacks and why hospitals need better security tools than just regular MFA. Stronger security measures such as endpoint security, network segmentation, and the adoption of updated and relevant security systems can help identify these threats before they escalate into full-scale breaches.

3.3. Case 3. A combined attack on the public sector

Government agencies faced a sophisticated attack combining both AiTM and BitB techniques [38]:

- **Initial Vector.** Targeted phishing campaign against government employees.
- **Attack Method.** Combined AiTM proxy and BitB popup techniques.
- **Impact.** Unauthorized access to sensitive government communications.
- **Scope.** Multiple agencies affected.
- **Detection.** Identified through pattern analysis of login attempts.
- **Key Vulnerability.** Limited implementation of advanced authentication protocols.

One serious case involved government employees getting fake emails that looked like real government messages. These emails led workers to a fake website that looked like their usual work portal. Even though the workers used their security codes (MFA) to try to stay safe, the attackers were able to steal both their passwords and security codes as they typed them in. This attack was very serious as it allowed malicious actors see private government files and messages.

Universities have also been attacked using these methods [39]. In one case, attackers sent fake emails that looked like they came from the university's computer

support team. The emails told students and teachers to fix problems with their Microsoft 365 accounts. Upon clicking the link, users encountered a fake Microsoft login window made using BitB techniques. Multiple users entered their credentials, which let the attackers get into their email accounts and see private university information. The university had to give everyone new passwords and make their security stronger after this happened.

Business Email Compromise (BEC) can also be conducted through the means of BitB. Attackers can infiltrate corporate email systems through phishing or credential theft, impersonating relevant higher-ups or officers to manipulate employees and users into initiating unauthorized money transfers. These attacks once again utilize social engineering in its scheme, exploiting human trust to have an easier entry point to technical vulnerabilities.

One notable instance involved the use of a company email account in an Alaska-based business [29]. A fraudulent wire transfer was initiated. A large sum of money was transferred through accounts, and the attackers ensured a smooth flow in the transfer by monitoring internal email activity. The attackers mimicked the communication style of company executives, making the process seem legitimate. The transaction was processed before anyone was able to detect the deception, causing significant financial loss [29].

These types of BitB attacks have grown to be more dangerous as they are able to bypass MFA. In some cases, attackers can use real-time session hijacking, where stolen credentials and MFA codes are immediately used to access accounts before they expire. This implies that even the most security-conscious users who trust and rely on MFA heavily are not fully protected. Organizations can often assume that MFA is a complete prevention method against phishing, but BitB attacks are an example that attackers have found ways to exploit human trust through the use of familiar login displays. These examples show that these modern phishing attacks can hurt any organization that has important information, not just banks and hospitals. Government agencies, universities, research institutions, and private companies must all develop more conscious and stronger security measures. Advanced detection and real-time are just a few ways in which BitB attacks can be combatted. Organizations must also perform frequent cybersecurity training and audits to ensure the safety of users, students, and employees.

3.4. Synthesis of the case studies

Analysis of these three cases reveals several important insights about modern phishing attacks. First, the researchers noticed that all attacks shared some basic features:

- Attackers initiated with fraudulent emails.
- Attackers breached outdated security measures that were believed to be dependable.

- The attacks were not noticed until significant damage had already been dealt. Looking deeper into these cases present us with three main security measures.

1. Security Technology Gaps – the obsolete ways of protecting systems are no longer sufficient for security [40]. Even organizations with good security found

themselves vulnerable to these new attacks. This shows the need for more advanced measures to detect and control these threats immediately.

2. The Human Factor – in each case, the attackers succeeded mainly because they deceived people, not because their technical flaws were exploited. This highlights how important it is to:

- a. Train employees regularly about new security threats;
- b. Help people understand what suspicious emails and websites look like;
- c. Create clear steps for reporting possible security problems.

3. Need for Better Monitoring – these cases show that organizations need to change how they watch for attacks. Instead of just checking if a user's login details are right, they need to watch for unusual behavior even after the user logs in [10]. Quick detection of abnormal activity could have prevented much of the damage in these cases.

4. Defense mechanisms

4.1. Detecting and preventing AiTM attacks

To stop AiTM attacks, companies need to use several different security tools and methods together [40]. One important approach is called Zero Trust security, which means checking every single request to access a system, even if it seems to come from someone inside the company. This helps catch attackers even if they manage to steal someone's login details.

Companies can also use special security systems that use Artificial Intelligence (AI) to spot unusual activity [10]. These systems can tell when something looks wrong, like when someone tries to log in from two different places at the same time. The AI can also find and block fake login pages before people enter their passwords.

Another way to improve security is to use something called certificate-based authentication. This means only devices with special digital certificates can access important systems [42]. Companies also use tools that watch how people log in and can quickly stop any suspicious activity.

Finally, companies should use tools that monitor their networks to spot signs of an AiTM attack, like unusual patterns in how data moves around or when someone tries to log in many times and fails. Using all these tools together, along with good security rules and training people about these risks, helps protect against AiTM attacks.

4.2. Mitigating BitB attacks

To protect against BitB attacks, organizations need to use both technical tools and user training [42]. Since these attacks create fake login windows that look real, it is important to teach people how to spot the differences. For example, real pop-up windows can be moved around and made bigger, while fake ones usually cannot. People should also learn to type website addresses themselves instead of clicking on links, which helps avoid these deceptive tactics.

Companies can use special security tools like anti-phishing browser add-ons and web filters to block dangerous websites [41]. These tools can spot suspicious things on websites, like fake login windows, and warn users before they enter their information. Some companies also use smart computer programs that can detect when someone's login behavior looks unusual.

One of the best ways to stay safe is to use special security devices called hardware security keys (like FIDO2 keys) [42]. These keys check both the user and the website during login, making sure that login information cannot be stolen by fake websites. Some companies also use something called browser isolation, which opens risky websites in a separate, protected space to keep users safe.

Finally, companies should regularly train their employees about these security risks and test them with fake phishing attempts. Using both security tools and good training helps protect people and companies from BitB attacks.

4.3. The role of AI in phishing defense

Artificial Intelligence (AI) has become very important in finding and stopping new types of phishing attacks like AiTM and BitB [40]. Old security tools used simple rules to spot attacks, but new phishing tactics keep emerging and are harder to catch. AI tools can look at lots of information quickly and learn to spot signs of attacks, even new ones they have not seen before.

One key way AI helps is by watching how people log into their accounts. AI can notice when something looks wrong – like when someone tries to log in from a strange place or device that they do not usually use [41]. When this happens, the AI can either stop the login or ask for extra proof that it is really the right person. AI is also good at spotting when someone might be using AiTM attacks by trying to log into many different accounts very quickly from different places.

For BitB attacks, AI looks at how login windows work and can tell when they are fake. The AI checks things like how the window moves and what it is made of, comparing it to how real login windows should work. This helps catch fake windows before people put in their passwords.

5. Conclusion

This research provides a comprehensive analysis of Adversary-in-The-Middle (AiTM) and Browser-in-the-Browser (BitB) attacks, detailing their mechanisms, how they bypass traditional security measures like Multi-Factor Authentication (MFA), and the increasing sophistication of these threats. Through case studies in banking, healthcare, and government institutions, we examine the real-world consequences of these attacks, including financial losses, data breaches, and service disruptions. Finally, we evaluate and propose effective defense strategies, such as Zero Trust architecture, AI-based detection, and employee training, to mitigate the risks posed by these evolving cyber threats.

Our investigation demonstrates that while MFA remains a cornerstone of modern cybersecurity, it is not impervious to advanced threats. AiTM and BitB attacks exploit inherent weaknesses in authentication frameworks, necessitating a

multi-layered approach to security. Organizations must implement supplementary protective measures such as encrypted communications, behavioral analytics, and continuous authentication monitoring to mitigate the risks posed by these evolving attack vectors. This study underscores the urgency of adapting cybersecurity strategies to counteract emerging threats, ensuring that authentication methods remain resilient against sophisticated adversaries.

For future study, the researchers recommend six key areas: expanding defense mechanisms beyond current architectures to include biometrics and blockchain verification; validating findings through practical attack simulations; examining regulatory frameworks' role in prevention; investigating human behavioral factors in phishing susceptibility; developing quantitative models for risk assessment; and studying the evolution of AiTM and BitB attacks, particularly with emerging AI-driven threats. These recommendations aim to strengthen cybersecurity defenses through comprehensive research and practical validation.

References

1. Denbigh-White, C. 2023 Verizon Data Breach Investigations Report: 7 Takeaways. Next DLP, 2024.
<https://www.nextdlp.com/resources/blog/seven-takeaways-from-2023-verizon-data-breach-investigations-report>
2. Ventura, J. Takeaways from the Verizon 2023 Data Breach Investigations Report. ThreatX, 2023.
<https://www.threatx.com/blog/takeaways-from-the-verizon-2023-data-breach-investigations-report/>
3. Bejamas. What Are AiTM Attacks and How to Protect Against Them. Descope, 2024.
<https://www.descope.com/learn/post/aitm-attack>
4. Brawner, M., K. Wojcieszek, G. Glass, R. Hicks. Rise in MFA Bypass Leads to Account Compromise. Kroll, 2023.
<https://www.kroll.com/en/insights/publications/cyber/mfa-bypass-leads-to-account-compromise>
5. Trivedi, A. Identifying Adversary-in-the-Middle (AiTM) Phishing Attacks through 3rd-Party Network Detection, 2023.
<https://techcommunity.microsoft.com/t5/microsoft-sentinel-blog/identifying-adversary-in-the-middle-aitm-phishing-attacks/ba-p/3991358>
6. Microsoft Threat Intelligence. Detecting and Mitigating a Multi-Stage AiTM Phishing and BEC Campaign. Microsoft Security Blog, 2023.
<https://www.microsoft.com/en-us/security/blog/2023/06/08/detecting-and-mitigating-a-multi-stage-aitm-phishing-and-bec-campaign/>
7. APWG. Phishing Activity Trends Report, 2024.
https://docs.apwg.org/reports/apwg_trends_report_q4_2023.pdf
8. Alkhalil, Z., C. Hewage, L. Nawaf, I. Khan. Phishing Attacks: A Recent Comprehensive Study and a New Anatomy. – Front. Comput. Sci., Vol. 3, March 2021. DOI: 10.3389/fcomp.2021.563060.
9. Arctic Wolf. History of Cybercrime. Arctic Wolf, 2024.
<https://arcticwolf.com/resources/blog/decade-of-cybercrime/>
10. Microsoft Threat Intelligence. Microsoft Digital Defense Report 2023 (MDDR), 2023.
<https://www.microsoft.com/en-us/security/security-insider/microsoft-digital-defense-report-2023>
11. Naqvi, B., K. Perova, A. Farooq, I. Makhdoom, S. Oyediji, J. Porras. Mitigation Strategies against the Phishing Attacks: A Systematic Literature Review. – Computers & Security, Vol. 132, 2023, 103387. DOI: 10.1016/j.cose.2023.103387.

12. IBM. Cost of a Data Breach 2024 | IBM. Cost of a Data Breach Report, 2024.
<https://www.ibm.com/reports/data-breach>
13. Desolda, G., L. Ferro, A. Marrella, M. Costabile, T. Catarci. Human Factors in Phishing Attacks: A Systematic Literature Review. – ACM Computing Surveys, Vol. **54**, 2022, No 35. DOI: 10.1145/3469886.
14. Birgisson, A., D. K. Smetters. So Long Passwords, Thanks for all the Phish. Google Online Security Blog, 2023.
<https://security.googleblog.com/2023/05/so-long-passwords-thanks-for-all-phish.html>
15. Proofpoint. 2024 State of the Phish Report: Phishing Statistics & Trends | Proofpoint US. Proofpoint, 2024.
<https://www.proofpoint.com/us/resources/threat-reports/state-of-phish>
16. NIST. NIST Special Publication 800-63B. Digital Identity Guidelines Authentication and Lifecycle Management, 2017.
<https://pages.nist.gov/sp800-63b.html>
17. Mandiant. M-Trends 2023 Special Report. Mandiant, 2023.
<https://www.mandiant.com/resources/reports/m-trends-2023-special-report>
18. MITRE. Modify Authentication Process, Technique T1556 – Enterprise | MITRE ATT&CK®. The MITRE Corporation, 2023.
<https://attack.mitre.org/techniques/T1556/>
19. Mohapatra, H., S. Rath, S. Panda, R. Kumar. Handling of Man-In-The-Middle Attack in WSN. – Intrusion Detection System, Vol. **8**, May 2020, pp. 1503-1510.
20. Amft, S., S. Hölterennhoff, N. Huaman, A. Krause, L. Simko, Y. Acar, S. Fahl. “We’ve Disabled MFA for You”: An Evaluation of the Security and Usability of Multi-Factor Authentication Recovery Deployments. – In: Proc. of 2023 ACM SIGSAC Conference on Computer and Communications Security (CCS ’23), 21 November 2023. Association for Computing Machinery, New York, NY, USA, 2023, pp. 3138-3152. DOI: 10.1145/3576915.3623180.
21. Gavazzi, A., R. Williams, E. Kirda, L. Lu, A. King, A. Davis, T. Leek. A Study of {Multi-Factor} and {Risk-Based} Authentication Availability. 2023, pp. 2043-2060.
<https://www.usenix.org/conference/usenixsecurity23/presentation/gavazzi>
22. Rajendran, H. H. Enhance MITM Attack Detection with Response Time in Secure Web Communication. Masters. Dublin, National College of Ireland, 2022.
<https://norma.ncirl.ie/6540/>
23. Chavoshi, H. R., A. H. Salasi, O. Payam, H. Khaloodeh. Man-in-the-Middle Attack Against a Network Control System: Practical Implementation and Detection. – In: Proc. of 64th IEEE International Scientific Conference on Information Technology and Management Science of Riga Technical University (ITMS), October 2023, pp. 1-6. DOI: 10.1109/ITMS59786.2023.10317671.
24. Cekerevac, Z., P. Cekerevac, L. Prigoda, F. Al-Naima. Security Risks from the Modern Man-in-the-Middle Attacks.
25. OWASP. OWASP Top 10, 2023: A10 Browser-in-the-Browser Attacks. Open Web Application Security Project, 2023.
<https://owasp.org/www-project-top-ten/>
26. Rescorla, E. Security Considerations for WebRTC. – Internet Engineering Task Force, 2021. DOI: 10.17487/RFC8826.
27. Perception Point. What Is a Browser-in-the-Browser (BitB) Attack? Perception Point, 2024.
<https://perception-point.io/guides/phishing/what-is-a-browser-in-the-browser-bitb-attack/>
28. Mozilla. Web Security Guidelines: Pop-up Authentication Windows. Mozilla Web Security, 2024.
https://infosec.mozilla.org/guidelines/web_security
29. FBI. Internet Crime Report 2023. Federal Bureau of Investigation.
<https://www.fbi.gov/contact-us/field-offices/sanfrancisco/news/fbi-releases-internet-crime-report>
30. FS-ISAC. Global Intelligence Office Report: BitB Attacks in Financial Services, 2024.
<https://www.fsisac.com/newsroom/fsisac-report-finds-global-cyberthreats-accelerate-as-cybercriminals-and-nation-state-actors-converge-and-collaborate>

31. Alsaffar, M., S. Aljaloud, B. A. Mohammed, Z. G. Al-Mekhlafi, T. S. Almurayziq, G. Alshammari, A. Alshammari. Detection of Web Cross-Site Scripting (XSS) Attacks. – Electronics, Vol. **11**, January 2022, No 14, 2212. DOI: 10.3390/electronics11142212.
32. Kusumo, W., A. Erlangga, M. R. Ramadhan. Potential Security Issues in Implementing IaaS and PaaS Cloud Service Models.
33. Europol. Internet Organised Crime Threat Assessment (IOCTA) 2023. Europol, 2024.
<https://www.europol.europa.eu/publication-events/main-reports/internet-organised-crime-threat-assessment-iocta-2023>
34. SWIFT. Swift Customer Security Controls Framework. Society for Worldwide Interbank Financial Telecommunication, 2024.
<https://www.swift.com/myswift/customer-security-programme-csp/security-controls>
35. FinCEN. Financial Trend Analysis: Cybercrime and Cyber-Enabled Crime Against Financial Institutions. Financial Crimes Enforcement Network, 2024.
<https://www.fincen.gov/resources/financial-trend-analyses>
36. HHS. Healthcare Cybersecurity Report: Rising Threats in Medical Systems, 2023.
<https://www.hhs.gov/about/news/2023/04/17/hhs-cybersecurity-task-force-provides-new-resources-help-address-rising-threat-cyberattacks-health-public-health-sector.html>
37. OCR. HIPAA Security Rule Compliance Guide, 2009.
<https://www.hhs.gov/hipaa/for-professionals/security/laws-regulations/index.html>
38. CISA. Federal Civilian Executive Branch Agency Cybersecurity Incident and Vulnerability Response Playbooks, 2024.
<https://www.cisa.gov/resources-tools/resources/federal-government-cybersecurity-incident-and-vulnerability-response-playbooks>
39. O'Brien, J. Higher Education Information Security Council Report: Phishing in Academia. EDUCAUSE Review, 2020.
<https://er.educause.edu/articles/2020/5/digital-ethics-in-higher-education-2020>
40. Gartner. Market Guide for Zero Trust Network Access. Gartner, 2023.
<https://www.gartner.com/en/documents/4632099>
41. Google. Safe Browsing: Protecting Web Users for 15 Years and Counting. Google Security Blog, 2023, 2024.
<https://www.googblog.com/category/online-security-blog/page/3/>
42. FIDO Alliance. FIDO2: Web Authentication (WebAuthn). FIDO Technical Specifications, 2023.
<https://fidoalliance.org/specs/fido-v2.2-rd-20230321/fido-client-to-authenticator-protocol-v2.2-rd-20230321.html>
43. BeEFProject. BeEF – The Browser Exploitation Framework Project.
<https://beefproject.com/>
44. OWASP. Cross Site Scripting (XSS) | OWASP Foundation.
<https://owasp.org/www-community/attacks/xss/>
45. Gillis, A. S. What is a Man-in-the-Browser Attack? Security.
<https://www.techtarget.com/searchsecurity/definition/man-in-the-browser>

Received: 12.12.2024, Revised Version: 21.02.2025, Accepted: 24.02.2025