

Chapter 1

Today's Cybersecurity Analyst

THE COMPTIA CYBERSECURITY ANALYST (CISA+) EXAM OBJECTIVES COVERED IN THIS CHAPTER INCLUDE:

✓ Domain 1.0: Security Operations

- 1.5 Explain the importance of efficiency and process improvement in security operations.
 - Standardize processes
 - Streamline operations
 - Technology and tool integration
- 1.6 Summarize concepts related to the use of AI in security operations.
 - AI risks
 - Governance
 - Use cases



Cybersecurity analysts are responsible for protecting the confidentiality, integrity, and availability of information and information systems used by their organizations. Fulfilling this responsibility begins with a commitment to a defense-in-depth approach to information security that uses multiple, overlapping security controls to achieve each cybersecurity objective. In Chapter 2, “System and Network Architecture,” we’ll discuss how to evolve that approach to incorporate a Zero Trust approach built on the principle that users inside or outside the network perimeter are never implicitly trusted.

In the first section of this chapter, you will learn how to assess the cybersecurity threats facing your organization and determine the risk that they pose to the confidentiality, integrity, and availability of your operations. In the sections that follow, you will learn about controls that you can put in place to secure networks and endpoints and evaluate the effectiveness of those controls over time.

Cybersecurity Objectives

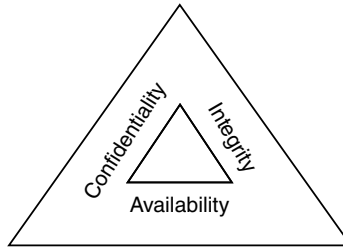
When most people think of cybersecurity, they imagine hackers trying to break into an organization’s system and steal sensitive information, ranging from Social Security and credit card numbers to top-secret military information. Although protecting sensitive information from unauthorized disclosure is certainly one element of a cybersecurity program, it is important to understand that information security actually has three complementary goals, as shown in Figure 1.1.

Confidentiality ensures that unauthorized individuals are not able to gain access to sensitive information. Cybersecurity professionals develop and implement security controls, including firewalls, access control lists, and encryption, to prevent unauthorized access to information. Attackers may seek to undermine confidentiality controls to achieve one of their objectives: the unauthorized disclosure of sensitive information.

Integrity ensures that there are no unauthorized modifications to information or systems, either intentionally or unintentionally. Integrity controls, such as hashing and integrity monitoring solutions, seek to enforce this requirement. Integrity threats may come from attackers seeking the alteration of information without authorization or non-malicious sources, such as a power spike causing the corruption of information.

Availability ensures that information and systems are ready to meet the needs of legitimate users at the time those users request them. Availability controls, such as fault

FIGURE 1.1 The three key goals of information security are confidentiality, integrity, and availability.



tolerance, clustering, and backups, seek to ensure that legitimate users may gain access as needed.

Similar to integrity threats, availability threats may come either from attackers seeking the disruption of access or from non-malicious sources, such as a fire destroying a datacenter that contains valuable information or services.

Cybersecurity analysts often refer to these three goals, known as the CIA Triad, when performing their work. They often characterize risks, attacks, and security controls as meeting one or more of the three CIA Triad goals when describing them.

Privacy vs. Security

Privacy and security are closely related concepts. We just discussed the three major goals of information security: confidentiality, integrity, and availability. These goals are all focused on the ways that an organization can protect its own data. Confidentiality protects data from unauthorized disclosure. Integrity protects data from unauthorized modification. Availability protects data from unauthorized denial of access.

Privacy controls have a different focus. Instead of focusing on ways that an organization can protect its own information, privacy focuses on the ways that an organization can use and securely share information that it has collected about individuals, in accordance with policy. This data, known as *personally identifiable information (PII)*, is often protected by regulatory standards and is always governed by ethical considerations. Organizations seek to protect the security of private information and may do so using the same security controls that they use to protect other categories of sensitive information, but privacy obligations extend beyond just security. Privacy extends to include the ways that an organization uses and shares the information that it collects and maintains with others.

Exam Note

Remember that privacy and security are complementary and overlapping, but they have different objectives.

The *Privacy Management Framework (PMF)* outlines 9 privacy objectives that organizations should strive to follow:

- **Management** says that the organization should document its privacy practices in a privacy policy and related documents.
- **Agreement, notice and communication** says that the organization should notify individuals about its privacy practices and inform individuals of the type of information that it collects and how that information is used.
- **Collection and creation** says that the organization should collect and create PII only for the purposes identified in the notice and consented to by the individual.
- **Use, retention, and disposal** says that the organization should only use information for identified purposes and must dispose of information properly when it is no longer needed.
- **Access** says that the organization should provide individuals with access to any information about that individual in the organization's records, and provide individuals with the opportunity to update and correct errors in their information.
- **Disclosure to third parties** says that the organization will disclose information to third parties only when consistent with notice and consent.
- **Security for privacy** says that PII will be protected against unauthorized access.
- **Data integrity and quality** says that the organization will maintain accurate and complete information.
- **Monitoring and enforcement** says that the organization will put business processes in place to ensure that it remains compliant with its privacy policy.

The PMF objectives are strong best practices for building a privacy program. In some jurisdictions and industries, privacy laws require the implementation of several of these objectives. For example, the European Union's (EU's) General Data Protection Regulation (GDPR) requires that organizations handling the data of individuals in the EU process personal information in a way that meets privacy requirements.

Evaluating Security Risks

Cybersecurity risk analysis is the cornerstone of any information security program. Analysts must take the time to thoroughly understand their own technology environments and the external threats that jeopardize their information security. A well-rounded cybersecurity risk assessment combines information about internal and external factors to help analysts understand the threats facing their organization and then design an appropriate set of controls to meet those threats.

Before diving into the world of risk assessment, we must begin with a common vocabulary. You must know three important terms to communicate clearly with other risk analysts: vulnerabilities, threats, and risks.

A *vulnerability* is a weakness in a device, system, application, or process that might allow an attack to take place. Vulnerabilities are internal technical or operational factors that may be controlled by cybersecurity professionals. For example, a web server that is running an outdated version of the Apache service may contain a vulnerability that would allow an attacker to conduct a denial-of-service (DoS) attack against the websites hosted on that server, jeopardizing their availability. Cybersecurity professionals within the organization have the ability to remediate this vulnerability by upgrading the Apache service to the most recent version that is not susceptible to the DoS attack.

A *threat* in the world of cybersecurity is an internal or external force (accidental, environmental, structural, or adversarial) that may exploit a vulnerability, resulting in an adverse outcome. For example, a hacker who would like to conduct a DoS attack against a website and knows about an Apache vulnerability poses a clear cybersecurity threat. Although many threats are malicious in nature, this is not necessarily the case. For example, an earthquake may also disrupt the availability of a website by damaging the datacenter containing the web servers. Earthquakes clearly do not have malicious intent. In most cases, cybersecurity professionals cannot do much to eliminate a threat. Hackers will hack and earthquakes will strike whether we like it or not.

A *risk* is the combination of a threat and a corresponding vulnerability that would result in an adverse outcome. Both a threat and a corresponding vulnerability must be present before a situation poses a risk to the security of an organization. For example, if a hacker targets an organization's web server with a DoS attack, but the server was patched so that it is not vulnerable to that attack, there is no risk because even though a threat is present (the hacker), there is no corresponding vulnerability. Similarly, a datacenter may be vulnerable to earthquakes because the walls are not built to withstand the extreme movements present during an earthquake, but it may be located in a region of the world where earthquakes do not occur. The datacenter may be vulnerable to earthquakes, but there is little to no threat of earthquake in its location, so the risk is negligible.

The relationship between risks, threats, and vulnerabilities is an important one, and it is often represented by this equation:

$$\text{Risk} = \text{Threat} \times \text{Vulnerability}$$

This is not meant to be a literal equation where you would actually plug in values. Instead, it is meant to demonstrate the fact that risks exist only when there is both a threat and a corresponding vulnerability that the threat might exploit. If either the threat or corresponding vulnerability is zero, the risk is also zero. Figure 1.2 shows this in another way: risks are the intersection of threats and vulnerabilities.

Organizations should routinely conduct risk assessments to take stock of their existing risk landscape. The National Institute of Standards and Technology (NIST) publishes a guide for conducting risk assessments that is widely used throughout the cybersecurity field as a foundation for risk assessments. The document is designated NIST Special Publication (SP) 800-30 rev. 1. Let's take one more look at how threats, vulnerabilities, and risks relate to each other, using the diagram shown in Figure 1.3, drawn from the NIST guide.

FIGURE 1.2 Risks exist at the intersection of threats and vulnerabilities. If either the threat or the corresponding vulnerability is missing, there is no risk.

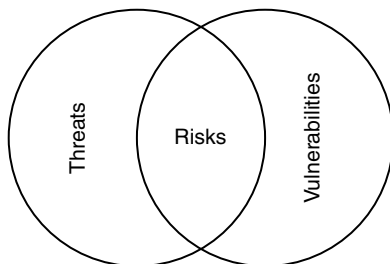
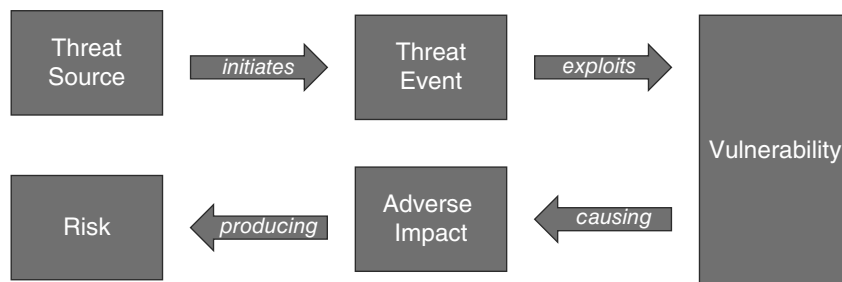


FIGURE 1.3 According to NIST, a threat source initiates a threat event, which exploits a vulnerability. That exploitation causes adverse impact, which produces organizational risk.



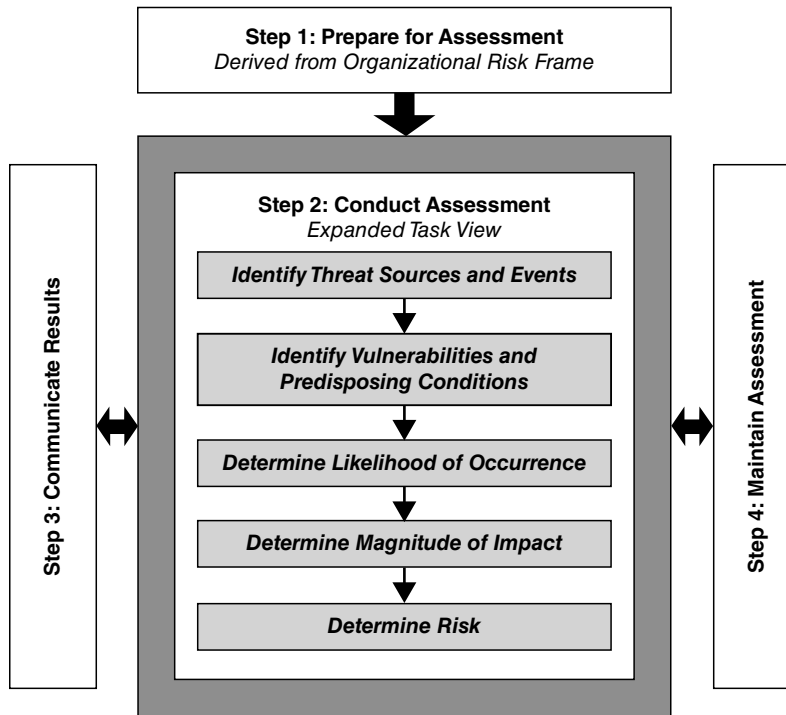
To identify and assess the risks facing an organization, NIST suggests using the risk assessment process shown in Figure 1.4.

Identify Threats

Organizations begin the risk assessment process by identifying the types of threats that exist in their threat environment. Although some threats, such as malware and spam, affect all organizations, other threats are targeted against specific types of organizations. For example, government-sponsored advanced persistent threat (APT) attackers typically target government agencies, military organizations, critical infrastructure, and companies that operate in related fields. It is unlikely that an APT attacker would target an elementary school.

NIST identifies four categories of threats that an organization might face and should consider in its threat identification process:

FIGURE 1.4 The NIST risk assessment process suggests that an organization should identify threats and vulnerabilities and then use that information to determine the level of risk posed by the combination of those threats and corresponding vulnerabilities.



SOURCE: Adapted from NIST SP 800-30 rev. 1

- **Adversarial threats** are individuals, groups, organizations, and nation-states that are attempting to deliberately undermine the security of an organization. Adversaries may include trusted insiders, competitors, suppliers, customers, business partners, or even nation-states. When evaluating an adversarial threat, cybersecurity analysts should consider the capability of the threat actor to engage in attacks, the intent of the threat actor, and the likelihood that the threat will target the organization.
- **Accidental threats** occur when individuals doing their routine work mistakenly perform an action that undermines security. For example, a system administrator might accidentally delete a critical disk volume, causing a loss of availability. When evaluating an accidental threat, cybersecurity analysts should consider the possible range of effects that the threat might have on the organization.
- **Structural threats** occur when equipment, software, or environmental controls fail due to the exhaustion of resources (such as running out of gas), exceeding their operational

capability (such as operating in extreme heat), or simply failing due to age. Structural threats may come from IT components (such as storage, servers, and network devices), environmental controls (such as power and cooling infrastructure), and software (such as operating systems and applications). When evaluating a structural threat, cybersecurity analysts should consider the possible range of effects that the threat might have on the organization.

- **Environmental threats** occur when natural or human-made disasters occur that are outside the control of the organization. These might include fires, flooding, severe storms, power failures, or widespread telecommunications disruptions. When evaluating environmental threats, cybersecurity analysts should consider common natural environmental threats to their geographic region, as well as how to appropriately prevent or counter human-made environmental threats.

The nature and scope of the threats in each of these categories will vary depending on the nature of the organization, the composition of its technology infrastructure, and many other situation-specific circumstances. That said, it may be helpful to obtain copies of the risk assessments performed by other, similar organizations as a starting point for an organization's own risk assessment or to use as a quality assessment check during various stages of the organization's assessment.

The Insider Threat

When performing a threat analysis, cybersecurity professionals must remember that threats come from both external and internal sources. In addition to the hackers, natural disasters, and other threats that begin outside the organization, rogue employees, disgruntled team members, and incompetent administrators also pose a significant threat to enterprise cybersecurity. As an organization designs controls, it must consider both internal and external threats, incorporating both a defense-in-depth and Zero Trust approach to cybersecurity.



NIST SP 800-30 rev. 1 provides a great deal of additional information to help organizations conduct risk assessments, including detailed tasks associated with each of these steps. This information is outside the scope of the Cybersecurity Analyst (CySA+) exam, but organizations preparing to conduct risk assessments should download and read the entire publication. It is available at <https://csrc.nist.gov/pubs/sp/800/30/r1/final>.

Identify Vulnerabilities

During the threat identification phase of a risk assessment, cybersecurity analysts focus on the external factors likely to impact an organization's security efforts. After completing threat identification, the focus of the assessment turns inward, identifying the vulnerabilities that those threats might exploit to compromise an organization's confidentiality, integrity, or availability.

Chapter 6, "Designing a Vulnerability Management Program," and Chapter 7, "Analyzing Vulnerability Scans," of this book focus extensively on the identification and management of vulnerabilities.

Determine Likelihood, Impact, and Risk

After identifying the threats and vulnerabilities facing an organization, risk assessors next seek out combinations of threat and vulnerability that pose a risk to the confidentiality, integrity, or availability of enterprise information and systems. This requires assessing both the likelihood that a risk will materialize and the impact that the risk will have on the organization if it does occur.

When determining the likelihood of a risk occurring, analysts should consider two factors. First, they should assess the likelihood that the threat source will initiate the risk. In the case of an adversarial threat source, this is the likelihood that the adversary will execute an attack against the organization. In the case of accidental, structural, or environmental threats, it is the likelihood that the threat will occur. The second factor that contributes is the likelihood that, if a risk occurs, it will actually have an adverse impact on the organization, given the state of the organization's security controls. After considering each of these criteria, risk assessors assign an overall likelihood rating. This may use categories, such as "low," "medium," and "high," to describe the likelihood qualitatively.

Risk assessors evaluate the impact of a risk using a similar rating scale. This evaluation should assume that a threat does take place and causes a risk to the organization and then attempt to identify the magnitude of the adverse impact that the risk will have on the organization. When evaluating this risk, it is helpful to refer to the three goals of cybersecurity shown in Figure 1.1, confidentiality, integrity, and availability, and then assess the impact that the risk would have on each of these goals.

Exam Note

The risk assessment process described here, using categories of "high," "medium," and "low," is an example of a qualitative risk assessment process. Risk assessments also may use quantitative techniques that numerically assess the likelihood and impact of risks. Quantitative risk assessments are beyond the scope of the CySA+ exam but are found on more advanced security exams, including the CompTIA SecurityX and Certified Information Systems Security Professional (CISSP) exams.

After assessing the likelihood and impact of a risk, risk assessors then combine those two evaluations to determine an overall risk rating. This may be as simple as using a matrix similar to the one shown in Figure 1.5 that describes how the organization assigns overall ratings to risks. For example, an organization might decide that the likelihood of a hacker attack is medium whereas the impact would be high. Looking this combination up in Figure 1.5 reveals that it should be considered a high overall risk. Similarly, if an organization assesses the likelihood of a flood as medium and the impact as low, a flood scenario would have an overall risk of low.

Reviewing Controls

Cybersecurity professionals use risk management strategies, such as risk acceptance, risk avoidance, risk mitigation, and risk transference, to reduce the likelihood and impact of risks identified during risk assessments. The most common way that organizations manage security risks is to develop sets of technical, physical, and administrative security controls that mitigate those risks to acceptable levels:

- *Technical controls* are systems, devices, software, and settings that work to enforce confidentiality, integrity, and/or availability requirements. Examples of technical controls include building a secure network and implementing endpoint security, two topics discussed later in this chapter.
- *Physical controls* are safeguards that prevent or deter unauthorized physical access to facilities, systems, or resources. Examples of physical controls include locks, security guards, fences, ID badges, surveillance cameras, and alarm systems.

FIGURE 1.5 Many organizations use a risk matrix to determine an overall risk rating based on likelihood and impact assessments.

Likelihood	High	Medium	High	Critical
	Medium	Low	Medium	High
	Low	Low	Low	Medium
		Low	Medium	High
		Impact		

- *Administrative controls* are practices and procedures that bolster cybersecurity. Examples of administrative controls include security policies, background checks, incident response plans, employee training, and acceptable use standards.

Building a Secure Network

Many threats to an organization's cybersecurity exploit vulnerabilities in the organization's network to gain initial access to systems and information. To help mitigate these risks, organizations should focus on building secure networks that keep attackers at bay. Examples of the controls that an organization may use to contribute to building a secure network include network access control (NAC) solutions; network perimeter security controls, such as firewalls; network segmentation; and the use of deception as a defensive measure.

Exam Note

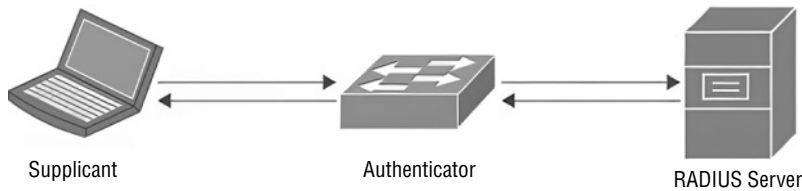
Much of the material in this chapter is not directly testable on the CySA+ exam. You won't, for example, find a question asking you how to configure a NAC solution. However, that doesn't mean that you don't need to know this material! You'll need to be familiar with the security controls in this chapter in order to analyze logs, recommend remediations, and conduct many other activities that *are* directly testable on the exam. The exam does assume that you have a basic familiarity with cybersecurity tools and techniques.

Network Access Control

One of the basic security objectives set forth by most organizations is controlling access to the organization's network. *Network access control (NAC)* solutions help security professionals achieve two cybersecurity objectives: limiting network access to authorized individuals and ensuring that systems accessing the organization's network meet basic security requirements.

The 802.1X protocol is a common standard used for NAC. When a new device wishes to gain access to a network, either by connecting to a wireless access point or plugging into a wired network port, the network challenges that device to authenticate using the 802.1X protocol. A special piece of software, known as a supplicant, resides on the device requesting to join the network. The supplicant communicates with a service known as the *authenticator* that runs on either the wireless access point or the network switch. The authenticator does not have the information necessary to validate the user itself, so it passes access requests along to an authentication server using the Remote Authentication Dial-In User Service (RADIUS) protocol. If the user correctly authenticates and is authorized to access the network, the switch or access point then joins the user to the network. If the user does

FIGURE 1.6 In an 802.1X system, the device attempting to join the network runs a NAC supplicant, which communicates with an authenticator on the network switch or wireless access point. The authenticator uses RADIUS to communicate with an authentication server.



not successfully complete this process, the device is denied access to the network or may be assigned to a special quarantine network for remediation. Figure 1.6 shows the devices involved in 802.1X authentication.

Many NAC solutions are available on the market, and they differ in two major ways:

Agent-Based vs. Agentless Agent-based solutions, such as 802.1X, require that the device requesting access to the network run special software designed to communicate with the NAC service. Agentless approaches to NAC conduct authentication in the web browser and do not require special software.

In-Band vs. Out-of-Band In-band (or inline) NAC solutions use dedicated appliances that sit in between devices and the resources that they wish to access. They deny or limit network access to devices that do not pass the NAC authentication process. The “captive portal” NAC solutions found in hotels that hijack all web requests until the guest enters a room number are examples of in-band NAC. Out-of-band NAC solutions leverage the existing network infrastructure and have network devices communicate with authentication servers and then reconfigure the network to grant or deny network access, as needed.

NAC solutions are often used simply to limit access to authorized users based on those users successfully authenticating, but they may also make network admission decisions based on other criteria. Some of the criteria used by NAC solutions are as follows:

Time of Day Users may be authorized to access the network only during specific time periods, such as during business hours.

Role Users may be assigned to particular network segments based on their role in the organization. For example, a college might assign faculty and staff to an administrative network that may access administrative systems while assigning students to an academic network that does not allow such access.

Location Users may be granted or denied access to network resources based on their physical location. For example, access to the datacenter network may be limited to systems physically present in the datacenter.

System Health NAC solutions may use agents running on devices to obtain configuration information from the device. Devices that fail to meet minimum security standards, such as having incorrectly configured host firewalls, outdated virus definitions, or missing security patches, may be either completely denied network access or placed on a special quarantine network where they are granted only the limited access required to update the system's security.

Administrators may create NAC rules that limit access based on any combination of these characteristics. NAC products provide the flexibility needed to implement the organization's specific security requirements for network admission.



You'll sometimes see the acronym NAC expanded to "Network Admission Control" instead of "network access control." In both cases, people are referring to the same general technology. Network Admission Control is a proprietary name used by Cisco for its network access control solutions.

Firewalls and Network Perimeter Security

NAC solutions are designed to manage the systems that connect directly to an organization's wired or wireless network. They provide excellent protection against intruders who seek to gain access to the organization's information resources by physically accessing a facility and connecting a device to the physical network. They don't provide protection against intruders seeking to gain access over a network connection. That's where firewalls enter the picture.

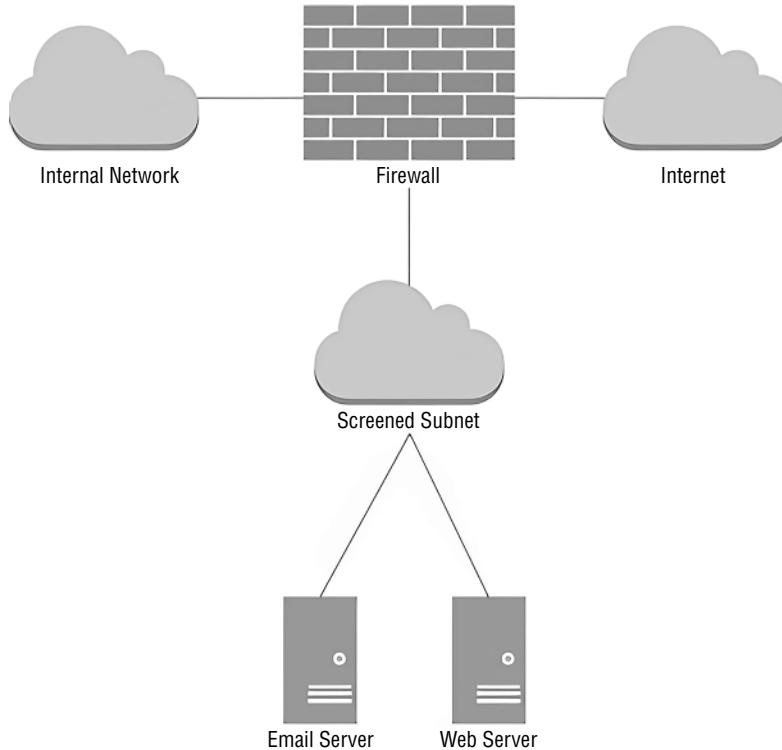
Network firewalls sit at the boundaries between networks and provide perimeter security. Much like a security guard might control the physical perimeter of a building, the network firewall controls the electronic perimeter. Firewalls are typically configured in the triple-homed fashion illustrated in Figure 1.7. Triple-homed simply means that the firewall connects to three different networks. The firewall in Figure 1.7 connects to the Internet, the internal network, and a special network zone known as the *screened subnet* (or demilitarized zone [DMZ]). Any traffic that wishes to pass from one zone to another, such as between the Internet and the internal network, must pass through the firewall.



Firewall capabilities can be extended to the cloud using firewall-as-a-service (FWaaS), cloud-native firewalls, and hybrid cloud solutions. This is a newer shift that enables dynamic, flexible, and scalable security for cloud environments, replacing the perimeter-based protection of traditional, on-premises hardware.

The screened subnet is a special network zone designed to house systems that receive connections from the outside world, such as web and email servers. Sound firewall designs place these systems on an isolated network where, if they become compromised, they pose

FIGURE 1.7 A triple-homed firewall connects to three different networks, typically an internal network, a screened subnet, and the Internet.



little threat to the internal network because connections between the screened subnet and the internal network must still pass through the firewall and are subject to its security policy.

Whenever the firewall receives a connection request, it evaluates it according to the firewall's rule base. This rule base is an access control list (ACL) that identifies the types of traffic permitted to pass through the firewall. The rules used by the firewall typically specify the source and destination IP addresses for traffic as well as the destination port corresponding to the authorized service. A list of common ports appears in Table 1.1. Firewalls follow the default deny principle, which says that if there is no rule explicitly allowing a connection, the firewall will deny that connection.



Some of the services listed in Table 1.1, particularly FTP, Telnet, HTTP, POP3, IMAP, SMTP, and LDAP, are insecure protocols that do not use encryption. You should generally avoid their use. However, it is important to know their ports so that you can properly write firewall rules to restrict these services, as desired.

TABLE 1.1 Common ports

Port	Service
20,21	FTP
22	SSH
23	Telnet
25	SMTP
53	DNS
67,68	DHCP
80	HTTP
110	POP3
123	NTP
161,162	SNMP
143	IMAP
389	LDAP
443	HTTPS
445	SMB
514	Syslog
587	SMTPS
636	LDAPS
993	IMAPS
995	POP3S
1433	SQL Server
1521	Oracle
3389	RDP

Several categories of firewalls are available on the market today, and they vary in both price and functionality:

- *Packet filtering firewalls* simply check the characteristics of each packet against the firewall rules without any additional intelligence. Packet filtering firewall capabilities are typically found in routers and other network devices and are very rudimentary firewalls.
- *Stateful inspection firewalls* go beyond packet filters and maintain information about session context and the state of each connection passing through the firewall. These are the most basic firewalls sold as stand-alone products.
- *Next-generation firewalls (NGFWs)* incorporate even more information into their decision-making process, including contextual information about users, applications, and business processes. They are the current state-of-the-art in network firewall protection and are quite expensive compared to stateful inspection devices.
- *Web application firewalls (WAFs)* are specialized firewalls designed to protect against web application attacks, such as SQL injection and cross-site scripting.
- *Proxy firewalls* act as an intermediary, or proxy, between an internal network and the Internet. They intercept network traffic at the application layer, providing enhanced security. Unlike traditional firewalls, which primarily filter based on IP addresses and ports, proxy firewalls deeply inspect application-level protocols, such as HTTP and FTP, to enforce policies, block malicious content, and conceal all internal IP addresses.

Network Segmentation

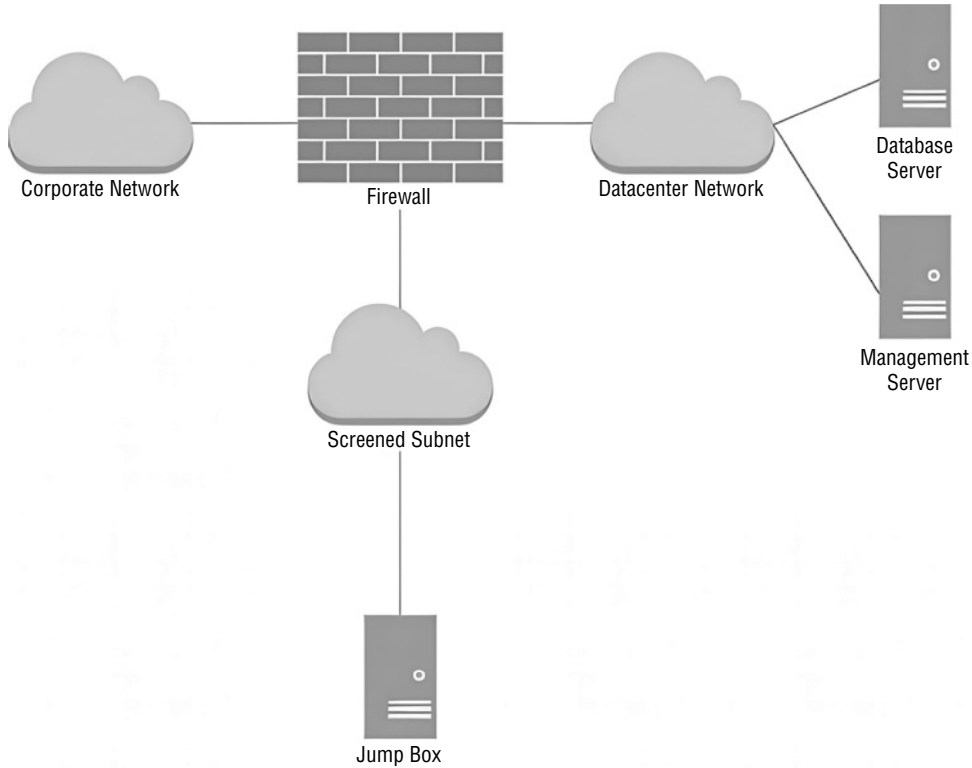
Firewalls use a principle known as *network segmentation* to separate networks of differing security levels from each other. This principle certainly applies to the example shown in Figure 1.7, where the internal network, screened subnet, and Internet all have differing security levels. The same principle may be applied to further segment the internal network into different zones of trust.

For example, imagine an organization that has several hundred employees and a large datacenter located in its corporate headquarters. The datacenter may house many sensitive systems, such as database servers that contain sensitive employee information, business plans, and other critical information assets. The corporate network may house employees, temporary contractors, visitors, and other people who aren't entirely trusted. In this common example, security professionals would want to segment the datacenter network so that it is not directly accessible by systems on the corporate network. This can be accomplished using a firewall, as shown in Figure 1.8.

The network shown in Figure 1.8 uses a triple-homed firewall, just as was used to control the network perimeter with the Internet in Figure 1.7. The concept is identical, except in this case the firewall is protecting the perimeter of the datacenter from the less trusted corporate network.

Notice that the network in Figure 1.8 also contains a screened subnet with a server called the *jump box*. The purpose of this server is to act as a secure transition point between the corporate network and the datacenter network, providing a trusted path between the two

FIGURE 1.8 A triple-homed firewall may also be used to isolate internal network segments of varying trust levels.



zones. System administrators who need to access the datacenter network should not connect devices directly to the datacenter network but should instead initiate an administrative connection to the jump box, using Secure Shell (SSH), the Remote Desktop Protocol (RDP), or a similar secure remote administration protocol. After successfully authenticating to the jump box, they may then connect from the jump box to the datacenter network, providing some isolation between their own systems and the datacenter network. Connections to the jump box should be carefully controlled and protected with strong multifactor authentication (MFA) technology.



MFA is a security mechanism that requires users to verify their identity using two or more authentication factors before gaining access to a system or service. This significantly enhances security by adding an extra layer of protection against unauthorized access, even if one factor is compromised. You'll learn more about MFA in Chapter 2, "System and Network Architecture."

Jump boxes may also be used to serve as a layer of insulation against systems that may only be partially trusted. For example, if you have contractors who bring equipment owned by their employer onto your network or employees bringing personally owned devices, you might use a jump box to prevent those systems from directly connecting to your company's systems.

Defense Through Deception

Cybersecurity professionals may wish to go beyond typical security controls and engage in active defensive measures that actually lure attackers to specific targets and seek to monitor their activity in a carefully controlled environment.

Honeypots are systems designed to appear to attackers as lucrative targets due to the services they run, vulnerabilities they contain, or sensitive information that they appear to host. The reality is that honeypots are designed by cybersecurity experts to falsely appear vulnerable and fool malicious individuals into attempting an attack against them. When an attacker tries to compromise a honeypot, the honeypot simulates a successful attack and then monitors the attacker's activity to learn more about their intentions. Honeypots may also be used to feed network *block lists*, blocking all inbound activity from any IP address that attacks the honeypot.

DNS sinkholes are a type of security tool that feeds false information to malicious software that infiltrates the enterprise network. When a compromised system attempts to obtain information from a DNS server about its command-and-control (C&C or C2) server, the DNS server detects the suspicious request and, instead of responding with the correct answer, responds with the IP address of a sinkhole system designed to detect and remediate the botnet-infected system.

Secure Endpoint Management

Laptop and desktop computers, tablets, smartphones, IoT devices, and other endpoint devices are a constant source of security threats on a network. These systems interact directly with end users and require careful configuration management to ensure that they remain secure and do not serve as the entry point for a security vulnerability on enterprise networks. Fortunately, by taking some simple security precautions, technology professionals can secure these devices against most attacks.

Hardening System Configurations

Operating systems are extremely complex pieces of software designed to perform thousands of different functions. The large code bases that make up modern operating systems are a frequent source of vulnerabilities, as evidenced by the frequent security patches issued by operating system vendors.

One of the most important ways that system administrators can protect endpoints is by hardening their configurations, making them as attack-resistant as possible. This includes updating their firmware, disabling any unnecessary services or ports on the endpoints to reduce their susceptibility to attack, ensuring that secure configuration settings exist on devices, and centrally controlling device security settings.

Patch Management

System administrators must maintain current security patch levels on all operating systems and applications under their care. Once the vendor releases a security patch, attackers are likely already aware of a vulnerability and may immediately begin preying on susceptible systems. The longer an organization waits to apply security patches, the more likely it becomes that they will fall victim to an attack. That said, enterprises should always test patches prior to deploying them on production systems and networks.

Fortunately, patch management software makes it easy to centrally distribute and monitor the patch level of systems throughout the enterprise. For example, Microsoft Intune allows administrators to quickly view the patch status of enterprise systems and remediate any systems with missing patches.

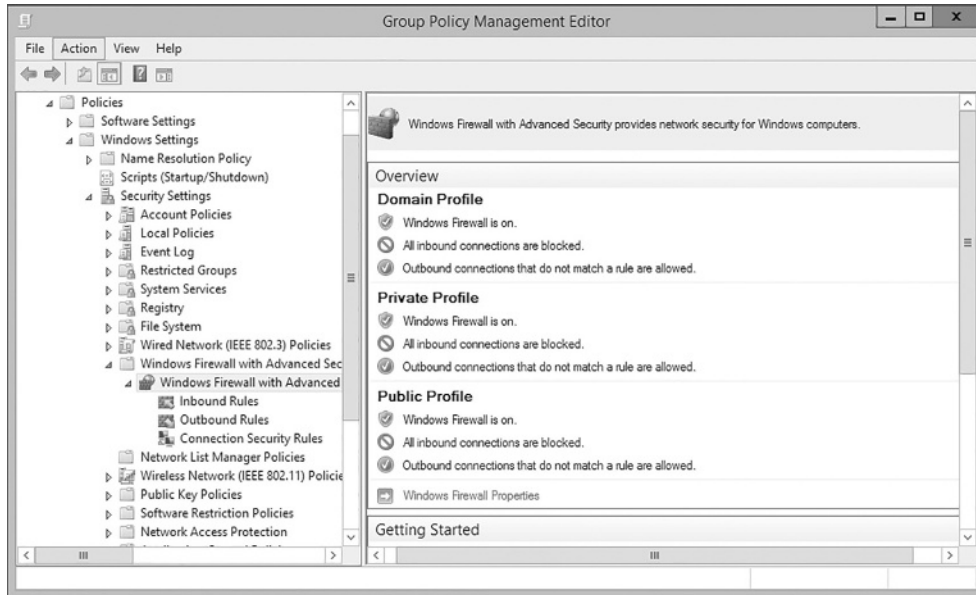
Compensating Controls

In some cases, security professionals may not be able to implement all of the desired security controls due to technical, operational, or financial reasons. For example, an organization may not be able to upgrade the operating system on retail point-of-sale (POS) terminals due to an incompatibility with the POS software. In these cases, security professionals should seek out compensating controls designed to provide a similar level of security using alternate means. In the POS example, administrators might place the POS terminals on a segmented, isolated network and use intrusion prevention systems to monitor network traffic for any attempt to exploit an unpatched vulnerability and block it from reaching the vulnerable host. This meets the same objective of protecting the POS terminal from compromise and serves as a compensating control.

Group Policies

Group Policies provide administrators with an efficient way to manage security and other system configuration settings across a large number of devices. Microsoft's Group Policy Object (GPO) mechanism allows administrators to define groups of security settings once and then apply those settings to either all systems in the enterprise or a group of systems based on role.

FIGURE 1.9 GPOs may be used to apply settings to many different systems at the same time.



For example, Figure 1.9 shows a GPO designed to enforce Windows Firewall settings on sensitive workstations. This GPO is configured to require the use of Windows Firewall and block all inbound connections.

Administrators may use GPOs to control a wide variety of Windows settings and create different policies that apply to different classes of systems.

Endpoint Security Software

Endpoint systems should also run specialized security software designed to enforce the organization's security objectives. At a minimum, this should include antivirus/antimalware software designed to scan the system for signs of malicious software that might jeopardize the security of the endpoint. Administrators may also choose to install host-based firewall software that serves as a basic firewall for that individual system, complementing network-based firewall controls or host intrusion prevention systems (HIPSs) that block suspicious network activity. Endpoint security software should report its status to a centralized management system that allows security administrators to monitor the entire enterprise from a single location.

Exam Note

Endpoint security goes beyond traditional antivirus/antimalware by incorporating endpoint detection and response (EDR) and extended detection and response (XDR) to continuously monitor for malicious activity, offer deeper visibility, and provide rapid threat containment and remediation.

Mandatory Access Controls

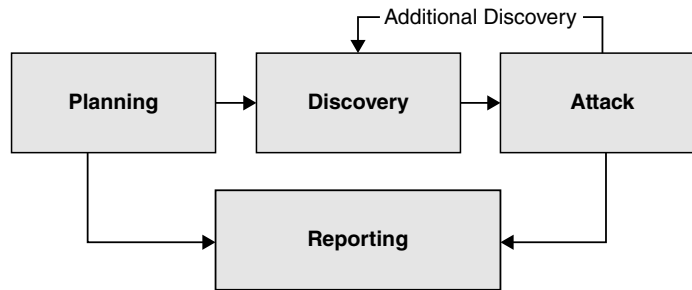
In highly secure environments, administrators may opt to implement a mandatory access control (MAC) approach to security. In a MAC system, administrators set all security permissions, and end users cannot modify those permissions. This stands in contrast to the discretionary access control (DAC) model found in most modern operating systems where the owner of a file or resource controls the permissions on that resource and can delegate them at their discretion.

MAC systems are very unwieldy and, therefore, are rarely used outside of very sensitive government and military applications. Security-Enhanced Linux (SELinux), an operating system developed by the U.S. National Security Agency, is an example of a system that enforces mandatory access controls.

Penetration Testing

In addition to bearing responsibility for the design and implementation of security controls, cybersecurity analysts are responsible for monitoring the ongoing effectiveness of those controls. *Penetration testing* is one of the techniques they use to fulfill this obligation. During a penetration test, the testers simulate an attack against the organization using the same information, tools, and techniques available to real attackers. They seek to gain access to systems and information and then report their findings to management. The results of penetration tests may be used to bolster an organization's security controls.

Penetration tests may be performed by an organization's internal staff or by external consultants. In the case of internal tests, they require highly skilled individuals and are quite time-consuming. External tests mitigate these concerns but are often quite expensive to conduct. Despite these barriers to penetration tests, organizations should try to perform them periodically since a well-designed and well-executed penetration test is one of the best measures of an organization's cybersecurity posture.

FIGURE 1.10 NIST divides penetration testing into four phases.

SOURCE: NIST SP 800-115 / National Institute of Standards and Technology / Public Domain

NIST divides penetration testing into the four phases shown in Figure 1.10.

Planning a Penetration Test

The planning phase of a penetration test lays the administrative groundwork for the test. No technical work is performed during the planning phase, but it is a critical component of any penetration test. There are three important rules of engagement to finalize during the planning phase:

Timing When will the test take place? Will technology staff be informed of the test? Can it be timed to have as little impact on business operations as possible?

Scope What is the agreed-on scope of the penetration test? Are any systems, networks, personnel, or business processes off-limits to the testers?

Authorization Who is authorizing the penetration test to take place? What should testers do if they are confronted by an employee or other individual who notices their suspicious activity?

These details are administrative in nature, but it is important to agree on them up front and in writing to avoid problems during and after the penetration test.



You should never conduct a penetration test without permission. Not only is an unauthorized test unethical, but it may be illegal.

Conducting Discovery

The technical work of the penetration test begins during the discovery phase when attackers conduct reconnaissance and gather as much information as possible about the targeted network, systems, users, and applications. This may include conducting reviews of publicly

available material, performing port scans of systems, using network vulnerability scanners and web application testers to probe for vulnerabilities, and performing other information gathering.



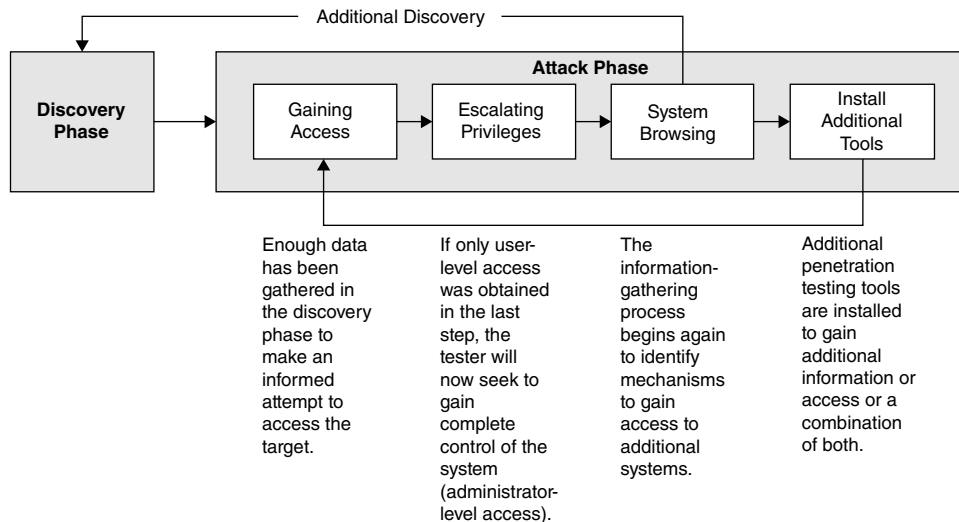
Vulnerability scanning is an important component of penetration testing. This topic is covered extensively in Chapters 6 and 7.

Executing a Penetration Test

During the attack phase, penetration testers seek to bypass the organization's security controls and gain access to systems and applications run by the organization. Testers often follow the NIST attack process shown in Figure 1.11.

In this process, attackers use the information gathered during the discovery phase to gain initial access to a system. Once they establish a foothold, they then seek to escalate their privileges until they gain complete administrative control of the system. From there, they can scan for additional systems on the network, install additional penetration testing tools, and begin the cycle anew, seeking to expand their footprint within the targeted organization. They continue this cycle until they exhaust the possibilities or the time allotted for the test expires.

FIGURE 1.11 The attack phase of a penetration test uses a cyclical process that gains a foothold and then uses it to expand access within the target organization.



SOURCE: NIST SP 800-115 / National Institute of Standards and Technology / Public Domain



If you're interested in penetration testing, CompTIA offers a companion certification to the CySA+ called PenTest+. The PenTest+ certification dives deeply into penetration testing topics.

Communicating Penetration Test Results

At the conclusion of the penetration test, the testers prepare a detailed report communicating the access they were able to achieve and the vulnerabilities they exploited to gain this access. The results of penetration tests are valuable security planning tools, because they describe the actual vulnerabilities that an attacker might exploit to gain access to a network. Penetration testing reports typically contain detailed appendixes that include the results of various tests and may be shared with system administrators responsible for remediating issues.

Training and Exercises

In addition to performing penetration tests, some organizations choose to run wargame exercises that pit teams of security professionals against one another in a cyber defense scenario. These exercises are typically performed in simulated environments, rather than on production networks, and seek to improve the skills of security professionals on both sides by exposing them to the tools and techniques used by attackers. Three teams are involved in most cybersecurity wargames:

- The red team plays the role of the attacker and uses reconnaissance and exploitation tools to attempt to gain access to the protected network. The red team's work is similar to that of the testers during a penetration test.
- The blue team is responsible for securing the targeted environment and keeping the red team out by building, maintaining, and monitoring a comprehensive set of security controls.
- The white team coordinates the exercise and serves as referees, arbitrating disputes between the team, maintaining the technical environment, and monitoring the results.

Cybersecurity wargames can be an effective way to educate security professionals on modern attack and defense tactics.

Efficiency and Process Improvement

Cybersecurity analysts perform a tremendous amount of routine work. From analyzing logs to assessing vulnerabilities, our work can be sometimes tedious and often error-prone. Strong cybersecurity teams invest time and money in improving their own processes, using standardization and automation to improve their efficiency, reduce the likelihood of error, and free up the valuable time of analysts for more significant work.

Standardize Processes

One of the first ways that teams can improve their efficiency is to create standardized processes for recurring activities. When you find yourself facing a task where you're inventing a solution as you go, that's a good sign that you might benefit from taking the time to create a standardized process for handling similar events in the future. This is especially true for activities you find yourself repeating time and time again.

Standardizing processes does more than reduce the effort required to react to a task; it also helps manage and facilitate team coordination. A documented process gives the team a shared language and a predictable rhythm. Analysts on different shifts can hand off a case using the same checkpoints and timelines rather than rewriting the plan in their own words. Roles and responsibilities become explicit so that work moves in parallel instead of piling up behind a single person. Service levels and decision points are defined in advance, which means new alerts drop into an established flow instead of starting an ad hoc debate.

In practice, this looks like turning to documentation for the scenario at hand and executing the same steps the team agreed upon yesterday, last week, and last quarter. The result is faster, more consistent outcomes and fewer misunderstandings, because every analyst knows what to do, when to do it, and who to do it with.

Cybersecurity Automation

Standardizing tasks also helps you identify opportunities for automation. You may be able to go beyond standardizing the work of team members and automate some responses to take people out of the loop entirely. *Security orchestration, automation, and response (SOAR)* platforms provide many opportunities to automate security tasks that cross between multiple systems.

A good first step is to work with your team to create an inventory of the activities you perform on a daily or weekly basis and identify those that are both repeatable and do not require human discretion. These are the best candidates for automation. For example, isolating an obviously compromised endpoint, running a reputation check against an indicator of compromise (IOC), or opening a ticket in a case management system may all be tasks that a SOAR workflow can complete automatically. Once these automations are in place, analysts still play a vital role in monitoring, refining, and expanding them, but the heavy lifting of execution is performed by the system.

SOAR platforms also offer opportunities to improve your organization's use of threat intelligence. By bringing information about emerging threats into your SOAR platform, you can enrich data about ongoing incidents and improve your ability to react to emerging cybersecurity situations. The SOAR platform provides you with the opportunity to combine information received through multiple threat feeds and develop a comprehensive picture of the cybersecurity landscape and your security posture.

Many security teams now use *infrastructure as code (IaC)* to manage their cloud and on-premises environments. With IaC, configuration and policy settings are stored in version-controlled code rather than being applied manually in consoles or GUIs. This approach allows organizations to deploy security controls in a consistent, repeatable manner

and to roll out changes to thousands of systems with a single update. For example, a team might use IaC templates to automatically enforce network segmentation rules, configure encryption defaults, or deploy logging agents across every new server spun up in the cloud.

Another important area for automation is rule and alert tuning. Analysts spend considerable time reviewing security alerts, many of which turn out to be false positives. Automating the enrichment of alert data, such as correlating with known threat intelligence feeds, adding user context, or checking whether an asset is already quarantined, reduces noise and ensures that high-value alerts rise to the top. Over time, security teams can also apply machine learning or simple conditional logic to automatically suppress alerts that match established benign patterns, further reducing analyst fatigue.

Automation also plays a role in how information is presented to decision makers. Well-designed *dashboards* give analysts, managers, and executives a real-time view of the organization's security posture. Automated dashboards can pull data from SIEMs, SOAR systems, vulnerability scanners, and threat intelligence feeds, presenting them in a unified view. This not only saves analysts the time of manually gathering data but also improves coordination across the team: everyone sees the same picture, using the same metrics, updated continuously. Dashboards can highlight key performance indicators (KPIs) such as mean time to detect (MTTD), mean time to respond (MTTR), and the volume of alerts handled, supporting both operational awareness and long-term process improvement.

Together, SOAR workflows, IaC deployments, automated alert tuning, and real-time dashboards demonstrate how standardization naturally leads to automation. Once processes are well defined, they can be encoded into tools and systems that reduce error, speed up response, and give teams the visibility they need to operate effectively in today's threat environment.

Technology and Tool Integration

Cybersecurity is full of tedious work. From performing vulnerability scans of large networks to conducting file integrity tests, we often need to use very repetitive processes to achieve our objectives. If we tried to get this work done manually, it would be so time-consuming that we would stop from exhaustion before we ever finished.

Workflow orchestration techniques help us automate many of these repetitive tasks. In addition to the SOAR platforms we've already discussed, there are two more ways that we can automate our workflows: *scripting*, which is writing code that automates our work, and *integration*, which uses vendor-provided interfaces to tie different products together.

Analysts commonly take advantage of *application programming interfaces (APIs)* as a primary means of integrating diverse security tools. APIs are programmatic interfaces to services that allow you to interact with that service without using web-based interfaces. You can normally perform the same actions with an API that you could perform at a service's web-based interface, but the API allows you to write code to automate those actions.

APIs are powerful to cybersecurity analysts because we can use them to reach into a wide variety of systems. Our security tools often offer APIs that we can leverage in our automation work, but so do many other technology services. We can use APIs to automate the provisioning of cloud resources, retrieve access logs from remote services, and automate many other routine tasks.

Webhooks allow us to send a signal from one application to another using a web request. For example, you might want to run a vulnerability scan every time your threat intelligence platform receives a report of a new vulnerability. In that case, you may be able to configure a webhook action in the threat intelligence platform that sends a request to the vulnerability scanner's API each time a new vulnerability is reported. That request could trigger the desired scan.

Plug-ins also provide an opportunity to increase integrations. These are small programs that run inside other programs, adding additional functionality. For example, you might use a browser plug-in to perform data enrichment. That plug-in might automatically pull up WHOIS and reputation data each time you hover over an IP address in your browser.

Exam Note

Be sure you are familiar with APIs, webhooks, and plug-ins as they are important cybersecurity integration tools, and they are specifically mentioned in the CySA+ exam objectives.

Bringing Efficiency to Incident Response

Incident response is one area of automation, as security teams bring the power of automation to what is often the most human-centric task in cybersecurity: investigating anomalous activity. While SOAR automation and other security tools may trigger an incident investigation, the work of the incident responder from that point forward has often been a very manual process that involves the application of tribal knowledge, personal experience, and instinct.

Enriching Incident Response Data

While incident response will likely always involve a significant component of human intervention, organizations have experienced success with automating portions of their incident response programs. One of the best starting points for incident response automation involves providing automated enrichment of incident response data to human analysts, saving them the tedious time of investigating routine details of an incident. For example, when an intrusion detection system identifies a potential attack, a security automation workflow can trigger a series of activities. These might include:

- Performing reconnaissance on the source address of the attack, including IP address ownership and geolocation information.
- Supplementing the initial report with other log information for the targeted system based upon a Security Information and Event Management (SIEM) query. SIEMs are an important security technology that you'll find covered in Chapter 3, "Malicious Activity."

- Triggering a vulnerability scan of the targeted system designed to assist in determining whether the attack has a high likelihood of success.

All of these actions can take place immediately upon the detection of the incident and appended to the incident report in the tracking system for review by a cybersecurity analyst. Teams seeking to implement incident response data enrichment will benefit from observing the routine activities of first responders and identifying any information gathering requirements that are possible candidates for automation.

Automate Portions of Incident Response Playbooks

In addition to enriching the data provided to cybersecurity analysts responding to an incident, automation may also play a role in improving the efficiency of response actions. Incident response teams commonly use *playbooks* that define the sequence of steps they will follow when responding to common incident types. Actions defined in a playbook may include activities designed to contain the incident, eradicate the effects of the incident from the network, and recover normal operations. You'll find more coverage of playbooks in Chapter 9, "Building an Incident Response Program." Teams may also refer to *runbooks* that IT teams use as references for performing routine procedures.

Exam Note

Be sure to keep the difference between playbooks and runbooks straight in your mind, as they are both covered by the CySA+ exam objectives. Playbooks are used to guide incident response efforts for a particular type of incident by directing the actions of the team. Runbooks provide step-by-step procedures for routine IT administrative tasks.

In some cases, it is reasonable to believe that the response to a security incident may be fully automated. For example, if a system on the local network begins emitting high volumes of traffic targeted at remote web servers, a reasonable analyst might conclude that the system is compromised and participating in a botnet-driven denial-of-service attack. The natural reaction of that analyst might then be to immediately quarantine the system to contain the damage and then send a case to the appropriate IT support technician to investigate and resolve the issue. That straightforward response is a prime candidate for automation. Integrations between the security automation platform, networking devices, and the ticketing systems can perform all these activities without any human intervention.

While the response to some simple incidents may be fully automated, many incidents have unique characteristics that require some degree of human intervention. Even in those cases, portions of the incident response playbook may still be automated beyond the data enrichment stage. For example, a cybersecurity analyst may have a dashboard that allows them to trigger a sequence of activities that blocks access to all resources for an individual user, quarantines suspect systems, or activates an incident escalation process.

Artificial Intelligence in Security Operations

Artificial intelligence (AI) is the use of computer systems to perform tasks that normally require human judgment, such as recognizing patterns, drawing conclusions, or making predictions. In cybersecurity, AI often relies on machine learning models that are trained with large volumes of data until they can identify anomalies, detect threats, or recommend actions with a level of speed and accuracy that would be impossible for a human analyst working alone.

The greatest strength of AI in security operations is its ability to handle data at scale. Security teams are flooded with logs, alerts, and network activity from dozens of different tools every day. Analysts cannot realistically examine all of this information directly, but AI-driven tools can analyze it continuously and surface the events that warrant closer attention. Instead of spending valuable time chasing down false positives, analysts can focus their energy on the alerts that matter most.

AI also offers adaptability that traditional signature-based or rule-based systems struggle to match. Rules and signatures must be updated manually and can quickly become outdated when attackers shift their techniques. Machine learning models, on the other hand, can be retrained on new data so that they evolve alongside the threat landscape. This ability to learn from fresh information makes AI especially valuable in identifying new or previously unseen attack methods.

Despite these advantages, AI is not a replacement for human expertise. Analysts still provide the context, reasoning, and oversight needed to make sound decisions about security events. AI can automate data processing, enrich alerts with additional context, and suggest likely classifications, but final responsibility for interpreting and responding to incidents remains with the human team. Used this way, AI becomes another layer of defense that strengthens security operations by reducing noise, accelerating response, and amplifying the effectiveness of skilled analysts.

AI Use Cases

Artificial intelligence is becoming a practical tool in daily security operations, helping analysts process information more efficiently and make better decisions under pressure. While AI is not a replacement for human expertise, it can take on many time-consuming tasks, provide enriched insights, and free analysts to focus on the highest-value work. The following examples highlight common ways that AI is applied in cybersecurity environments today:

- **Comparing artifacts:** AI can quickly compare files, registry entries, or network captures to identify similarities and differences that may indicate malicious activity or data manipulation.

- **Analyzing log files:** AI models excel at parsing large and complex log sets, spotting patterns, and flagging anomalies that would be difficult to detect with manual review.
- **Document creation:** Generative AI tools can draft reports, incident summaries, or compliance documentation, giving analysts a starting point that can be refined with human review.
- **Incident investigation:** AI can enrich alerts with threat intelligence, highlight probable root causes, and suggest next steps, helping analysts move from detection to resolution more quickly.
- **Event correlation:** By linking related alerts and signals across multiple systems, AI can reduce alert fatigue and present a clearer picture of an unfolding attack campaign.
- **Automation and orchestration:** AI-driven workflows can trigger automated responses, such as isolating endpoints or blocking IP addresses, and coordinate actions across different tools to streamline security operations.

EXAM NOTE

On the exam, expect questions that test your understanding of where AI can be applied in security operations. You should be able to distinguish between tasks suited for AI, such as analyzing log files or event correlation, and those that still require human judgment, such as deciding whether to notify law enforcement after an incident.

AI Governance

As organizations adopt artificial intelligence in their security operations, they must also establish governance practices that ensure its safe and responsible use. AI offers significant benefits, but without clear oversight, it can create new risks ranging from regulatory violations to unintentional misuse. *Governance* frameworks provide the structure for balancing innovation with accountability, both internally and externally.

One key consideration is legal or regulatory compliance. Many industries are subject to strict rules regarding the handling of sensitive information, and these requirements apply just as much to AI systems as traditional security tools. For example, if an organization uses AI to process PII, it must ensure that the system complies with privacy laws such as the GDPR or the Health Insurance Portability and Accountability Act (HIPAA). In the European context, organizations must also consider the emerging EU AI Act, which classifies AI systems according to risk and imposes obligations on high-risk applications. Security automation tools that make consequential decisions (such as isolating systems or blocking user accounts) may fall under the act's high-risk category and thus demand additional controls and governance. Failing to apply compliance checks to AI can expose organizations to legal liability, regulatory penalties, and reputational damage.

Equally important are AI usage policies. A usage policy defines how AI tools may and may not be used within the organization. These policies should clarify which tasks AI is authorized to perform, what types of data it may process, and the level of human oversight required. For example, a policy might allow AI to enrich incident data or draft reports but require human review before any automated system takes actions that could disrupt business operations, such as terminating sessions, blocking access, or shutting down servers. Policies should also address data security, prohibiting the use of public AI tools to process confidential or regulated information unless explicit safeguards are in place. They should account for evolving regulatory standards like the EU AI Act and ensure that the organization's AI systems remain compliant across jurisdictions.

Strong governance combines compliance and policy with education and accountability. Analysts need to be trained on the organization's AI usage policies: understanding not just *how* to use AI tools, but also *why* certain AI actions require review. Leadership must periodically audit AI-enabled processes and systems to ensure alignment with legal requirements, documentation obligations, and risk categories under frameworks such as the EU AI Act. Systems themselves should produce audit logs that record AI decisions, data sources, and human overrides. By implementing these layered governance practices, organizations can harness the power of AI in security operations while minimizing risk, protecting stakeholder rights, and maintaining trust across global regulatory environments.

AI Risks

Although artificial intelligence can bring powerful benefits to security operations, it also introduces new risks that analysts must understand and manage. These risks may come from the way AI systems process information, the data used to train or prompt them, or the possibility of attackers deliberately manipulating AI tools. Recognizing these risks allows organizations to implement safeguards that keep AI a reliable asset rather than a liability.

- *Hallucinations*: AI systems sometimes generate outputs that are incorrect, misleading, or entirely fabricated. In security operations, a hallucination could mean an AI tool misclassifies benign activity as malicious or invents evidence during an investigation, leading to wasted effort or inappropriate action.
- *Data exposure*: When sensitive information is provided to an AI system, particularly one hosted by a third party, there is a risk that the data may be stored, reused, or inadvertently exposed. This is a particular concern if analysts paste logs or artifacts containing personally identifiable information (PII) or confidential business data into public AI tools.
- *Model poisoning*: Attackers may attempt to corrupt the training data or the feedback loops that shape an AI system's behavior. By introducing malicious inputs, they can skew results, causing the model to overlook true threats or classify malicious activity as safe.
- *Malicious prompts*: Generative AI systems can be manipulated through carefully crafted inputs, called *prompt injection* attacks. A malicious prompt could override safeguards, trick the AI into revealing sensitive information, or cause it to take harmful automated actions.

AI risk management requires vigilance from both security teams and organizational leadership. Analysts must remain skeptical of AI-generated outputs, ensuring human oversight before taking critical actions. Organizations should implement technical safeguards such as data sanitization, access controls, and model monitoring, along with clear policies limiting the use of public AI services for sensitive data. By anticipating these risks and embedding safeguards into daily workflows, security teams can make effective use of AI while protecting against its potential downsides.

Exam Essentials

Know the three goals of cybersecurity. *Confidentiality* ensures that unauthorized individuals are not able to gain access to sensitive information. *Integrity* ensures that there are no unauthorized modifications to information or systems, either intentionally or unintentionally. *Availability* ensures that information and systems are ready to meet the needs of legitimate users at the time those users request them.

Understand how networks are made more secure through the use of network access control, firewalls, and segmentation. Network access control (NAC) solutions help security professionals achieve two cybersecurity objectives: limiting network access to authorized individuals and ensuring that systems accessing the organization's network meet basic security requirements. Network firewalls sit at the boundaries between networks and provide perimeter security. Network segmentation uses isolation to separate networks of differing security levels from each other.

Understand how endpoints are made more secure through the use of hardened configurations, patch management, Group Policy, and endpoint security software. Hardening configurations include disabling any unnecessary services on the endpoints to reduce their susceptibility to attack, ensuring that secure configuration settings exist on devices, and centrally controlling device security settings. Patch management ensures that operating systems and applications are not susceptible to known vulnerabilities. Group Policy allows the application of security settings to many devices simultaneously, and endpoint security software protects against malicious software and other threats.

Understand how automation and orchestration streamline repetitive security tasks. Automation removes repetitive, low-value tasks from analysts, while orchestration coordinates workflows across tools. Security Orchestration, Automation, and Response (SOAR) platforms can isolate compromised endpoints, enrich incident data, or open tickets without human intervention. Infrastructure as code (IaC) applies the same principles to system configuration, allowing organizations to enforce consistent security policies at scale.

Know how data enrichment improves alert quality and response times. Analysts spend much of their time reviewing alerts, many of which are false positives. Automating data

enrichment, such as correlating alerts with threat intelligence, adding user context, or checking system health, helps highlight high-value alerts. Over time, rules and machine learning can automatically suppress known benign activity, reducing noise and analyst fatigue.

Be able to explain the role of dashboards in improving visibility and decision-making. Dashboards provide real-time views of security posture by pulling data from SIEMs, SOAR platforms, and other tools into one place. They display key metrics like mean time to detect (MTTD) and mean time to respond (MTTR), helping analysts and executives understand current threats and long-term trends. Automated dashboards improve coordination by ensuring all stakeholders see the same up-to-date information.

Know how APIs, webhooks, and plug-ins enable tool integration. APIs allow security teams to programmatically interact with systems, automating tasks such as retrieving logs or provisioning cloud resources. Webhooks enable one tool to trigger actions in another, such as running a scan when a new vulnerability is discovered. Plug-ins extend the functionality of existing tools, often adding enrichment or specialized automation features. Together, these integrations reduce manual effort and increase efficiency.

Understand the role of playbooks and runbooks in incident response efficiency. Playbooks provide predefined steps for handling common incident types, ensuring consistency across responses. Portions of these playbooks can be automated, such as quarantining a compromised system or revoking user access, to speed containment. This blend of automation and human judgment creates faster and more accurate incident handling. Runbooks, on the other hand, provide step-by-step procedures for performing routine IT administrative tasks.

Know the potential risks of using AI in security operations. AI introduces risks such as hallucinations, where models produce convincing but false information, and data exposure when sensitive logs or artifacts are entered into AI systems. Models can also be poisoned by malicious training data or manipulated with prompt injection. These risks mean human oversight and proper safeguards are always required.

Be able to summarize AI use cases in security operations. AI can help analysts compare suspicious artifacts, analyze massive log sets, and generate draft reports. It can enrich incidents with threat intelligence, correlate related events, and automate routine tasks. These applications reduce analyst workload and accelerate response but should complement, not replace, human decision-making.

Understand the need for governance and compliance in AI usage. AI use must align with legal and regulatory requirements such as GDPR, HIPAA, or the EU AI Act. Organizations should create AI usage policies that define acceptable use, data handling practices, and required oversight. Governance also includes auditing AI outputs, training staff on proper use, and maintaining logs of AI-driven decisions to ensure accountability and trust.

Lab Exercises

Activity 1.1: Create an Inbound Firewall Rule

In this lab, you will verify that the Windows Defender Firewall is enabled on a server and then create an inbound firewall rule that blocks file and printer sharing.

These lab instructions were written to run on a system running Windows Server 2022. The process for working on other versions of Windows Server is quite similar, although the exact names of services, options, and icons may differ slightly.



You should perform this lab on a test system. Enabling file and printer sharing on a production system may have undesired consequences. The easiest way to get access to a Windows Server 2022 system is to create an inexpensive cloud instance through Amazon Web Services (AWS) or Microsoft Azure.

Part 1: Verify that Windows Defender Firewall is enabled

1. Open Control Panel for your Windows Server.
2. Choose System and Security.
3. Under Windows Defender Firewall, click Check Firewall Status.
4. Verify that the Windows Defender Firewall state is set to On for Private networks. If it is not on, enable the firewall by using the “Turn Windows Defender Firewall on or off” link on the left side of the window.

Part 2: Create an inbound firewall rule that allows file and printer sharing

1. On the left side of the Windows Defender Firewall Control Panel, click “Allow an app or feature through Windows Defender Firewall.”
2. Scroll down the list of applications and find File and Printer Sharing.
3. Check the box to the left of that entry to block connections related to File and Printer Sharing.
4. Confirm that the Private box to the right of that option was automatically selected. This allows File and Printer Sharing only for other systems on the same local network. The box for public access should be unchecked, specifying that remote systems are not able to access this feature.
5. Click OK to apply the setting.

Activity 1.2: Create a Group Policy Object

In this lab, you will create a Group Policy Object and edit its contents to enforce an organization's password policy.

These lab instructions were written to run on a system running Windows Server 2022. The process for working on other versions of Windows Server is quite similar, although the exact names of services, options, and icons may differ slightly. To complete this lab, your Windows Server must be configured as a domain controller.

1. Open the Group Policy Management application. (If you do not find this application on your Windows Server, it is likely that it is not configured as a domain controller.)
2. Expand the folder corresponding to your Active Directory forest.
3. Expand the Domains folder.
4. Expand the folder corresponding to your domain.
5. Right-click the Group Policy Objects folder and select New from the pop-up menu.
6. Name your new GPO **Password Policy** and click OK.
7. Click the Group Policy Objects folder.
8. Right-click the new Password Policy GPO and select Edit from the pop-up menu.
9. When Group Policy Management Editor opens, expand the Policies folder under the Computer Configuration section.
10. Expand the Windows Settings folder.
11. Expand the Security Settings folder.
12. Expand the Account Policies folder.
13. Click Password Policy.
14. Double-click Maximum Password Age.
15. In the pop-up window, select the Define This Policy Setting check box and set the expiration value to 90 days.
16. Click OK to close the window.
17. Click OK to accept the suggested change to the minimum password age.
18. Double-click the Minimum Password Length option.
19. As in the prior step, click the box to define the policy setting and set the minimum password length to 12 characters.
20. Click OK to close the window.

21. Double-click the Password Must Meet Complexity Requirements option.
22. Click the box to define the policy setting and change the value to Enabled.
23. Click OK to close the window.
24. Click the X to exit Group Policy Management Editor.

You have now successfully created a Group Policy Object that enforces the organization's password policy. You can apply this GPO to users and/or groups as needed.

Activity 1.3: Write a Penetration Testing Plan

For this activity, you will design a penetration testing plan for a test against an organization of your choosing. If you are employed, you may choose to use your employer's network. If you are a student, you may choose to create a plan for a penetration test of your school. Otherwise, you may choose any organization, real or fictitious, of your choice.

Your penetration testing plan should cover the three main criteria required before initiating any penetration test:

- Timing
- Scope
- Authorization

One word of warning: You should not conduct a penetration test without permission of the network owner. This assignment only asks you to design the test on paper.

Activity 1.4: Recognize AI Risks

Match each of the AI risks in this table with the correct description.

Hallucinations	An attacker attempts to corrupt the training data or the feedback loops that shape an AI system's behavior.
Data exposure	Generative AI systems can be manipulated through carefully crafted inputs, called prompt injection attacks.
Model poisoning	An AI tool misclassifies benign activity as malicious or invents evidence during an investigation.
Malicious prompts	An analyst pastes logs or artifacts containing personally identifiable information (PII) or confidential business data into public AI tools.

Correct Answers

Hallucinations	An AI tool misclassifies benign activity as malicious or invents evidence during an investigation.
Data exposure	An analyst pastes logs or artifacts containing personally identifiable information (PII) or confidential business data into public AI tools.
Model poisoning	An attacker attempts to corrupt the training data or the feedback loops that shape an AI system's behavior.
Malicious prompts	Generative AI systems can be manipulated through carefully crafted inputs, called prompt injection attacks.

Activity 1.5: Recognize Processes, Operations, and Technology Tools

Match each of the processes, operations, and technology tools in this table with the correct description.

Playbook	Used by IT operations for reference for routine procedures administrators perform.
Runbook	Platform that offers opportunities to improve your organization's use of threat intelligence. By bringing information about emerging threats into this platform, you can enrich data about ongoing incidents and improve your ability to react to emerging cybersecurity situations.
SOAR	Gives security analysts, managers, and executives a real-time view of the organization's security posture. These dashboards can pull data from SIEMs, SOAR systems, vulnerability scanners, and threat intelligence feeds, presenting them in a unified view.
IaC	Provide an opportunity to increase integrations. These are small programs that run inside other programs, adding additional functionality.
Dashboard creation	Are a primary means of integrating diverse security tools. These are programmatic interfaces to services that allow you to interact with that service without using web-based interfaces.

APIs	Allow us to send a signal from one application to another using a web request.
Webhooks	With this, configuration and policy settings are stored in version-controlled code rather than being applied manually in consoles or GUIs. This approach allows organizations to deploy security controls in a consistent, repeatable manner and to roll out changes to thousands of systems with a single update.
Pug-ins	Provides manual orchestration of incident response. For example, specific incidents and threats each have their own.

Review Questions

1. A security team finds that each analyst handles malware incidents in a completely different way, leading to confusion and inconsistent results. What process improvement should the team implement to ensure future incidents are handled more uniformly and efficiently?
 - A. Continue allowing each analyst to use their personal approach.
 - B. Develop a standardized playbook for malware incident response.
 - C. Change job descriptions so that the same person handles all malware incidents.
 - D. Delay responding to malware incidents until a manager can advise.
2. Attackers have started using new techniques that don't match any known signatures or rules, and the current security tools are missing these attacks. How could artificial intelligence-based tools help address this gap?
 - A. AI will force attackers to revert to old techniques that have signatures.
 - B. AI operates on fixed rules, so it wouldn't help with unknown threats.
 - C. AI would ignore any activity that isn't in its initial database of threats.
 - D. AI models can learn from new data and adapt to novel attack patterns more quickly than static, signature-based systems.
3. During a critical incident, the incident response team wasted time debating responsibilities and when to escalate decisions. What should they have in place beforehand to avoid these delays?
 - A. A documented, standardized incident response plan defining roles and escalation points
 - B. An expectation that team members will sort out roles during the incident
 - C. A policy of waiting for management guidance in every incident
 - D. A larger team to improve the diversity of skills
4. Which of the following security tasks is most appropriate to hand over to an AI-based system rather than relying on a human analyst?
 - A. Writing incident postmortem reports with recommendations for company culture changes
 - B. Defining acceptable risk tolerance levels for the organization
 - C. Correlating thousands of alerts from multiple tools to identify an attack pattern
 - D. Writing customized firewall policies tailored to evolving business needs

5. After formalizing a consistent process for investigating malware alerts, the security team notices many steps are repetitive and do not require judgment. What is the next step they should consider to further improve efficiency?
 - A. Abandon the standardized process since it revealed unnecessary steps.
 - B. Assign an intern to perform all the repetitive tasks manually.
 - C. Let analysts continue performing the repetitive steps to maintain control.
 - D. Automate those repetitive steps using a security orchestration and automation tool.
6. The security team has a rule that no one may paste sensitive internal logs or confidential data into any public AI or cloud-based AI tool. Which risk is this policy designed to mitigate?
 - A. Malicious prompts
 - B. Model poisoning
 - C. Hallucinations
 - D. Data exposure
7. Which of the following security tasks is the best candidate for automation using scripts or a SOAR platform?
 - A. Brainstorming new threat scenarios for the upcoming quarter
 - B. Conducting a one-on-one interview with an employee about a security policy violation
 - C. Automatically isolating a workstation when it is confirmed as infected with malware
 - D. Developing a custom strategy for a zero-day cyberattack
8. Which one of the following laws or regulations governs the use of AI to conduct security-related work in the European Union?
 - A. AI Act
 - B. GDPR
 - C. PCI DSS
 - D. HIPAA
9. A security analyst integrates multiple threat intelligence feeds into the team's automated response system. What is the primary benefit of this integration?
 - A. Threat actors are automatically identified and blocked across all systems in real time.
 - B. The system's firewall rules are rewritten dynamically to prevent future attacks.
 - C. The integration reduces the need for endpoint protection software.
 - D. Alerts are enriched with current threat context, allowing analysts to prioritize real threats more effectively.

10. A security team is working on an integration where its threat intelligence platform will trigger a vulnerability scan automatically when it receives word of a new critical CVE. Which integration method fits this workflow best?
- A. Webhook call from the intel platform to the scanner's API
 - B. Manual export of the CVE list to a CSV for the scanner
 - C. Nightly email to the SOC with new CVEs
 - D. Local script that parses a weekly PDF advisory
11. A cloud operations team needs to enforce that every new virtual server launched has standard security settings (such as firewall rules and logging) applied automatically. Which process improvement will best help them achieve this?
- A. Infrastructure as code (IaC)
 - B. Security information and event management (SIEM)
 - C. Security orchestration, automation, and response (SOAR)
 - D. Threat hunting
12. The security team finds that an AI tool occasionally flags benign activity as malicious due to lack of context. When this happens, the team double-checks the AI's findings before taking action. What is the best practice being demonstrated here?
- A. Validating AI output
 - B. Trusting the model's judgment
 - C. Fine-tuning the model to improve future results
 - D. Removing humans from the loop to increase efficiency
13. Over the past month, analysts have been flooded by false-positive security alerts. Many alerts are irrelevant, causing fatigue and wasted effort. What should the team do to streamline their operations and reduce this noise?
- A. Increase the sensitivity of intrusion detection systems to catch more threats.
 - B. Implement automated alert tuning to enrich alert data and filter out known benign events.
 - C. Hire additional analysts to manually review every alert.
 - D. Replace all detection tools with a single vendor solution.
14. A cloud team wants every new VM to have logging agents and baseline firewall rules without manual steps. Which approach aligns with tool integration and repeatability?
- A. Create a shared spreadsheet with setup steps for engineers to follow.
 - B. Open a ticket for each VM so IT Operations can apply settings.
 - C. Let each engineer script their own one-off bootstrap sequence.
 - D. Use infrastructure as code templates that call provider APIs during provisioning.

15. The SOC has written rules to automatically suppress alerts that match patterns of benign activity and has even trained a model to recognize known safe behavior. What is the expected result of this effort?
- A. Fewer false-positive alerts, allowing analysts to concentrate on genuinely suspicious events.
 - B. False negatives will no longer occur in the environment.
 - C. A higher number of alerts since the system is monitoring more conditions.
 - D. The SOC will no longer need to maintain detection rules manually.
16. An analyst hovers over an IP address in a browser and sees WHOIS and reputation details appear. Which integration use case does this describe?
- A. Event correlation
 - B. Log analysis
 - C. Tool orchestration
 - D. Data enrichment
17. The chief information security officer (CISO) wants a live view of key security metrics like mean time to detect (MTTD), mean time to respond (MTTR), and the number of alerts handled per day, all in one place. What solution should the team implement?
- A. Print out a daily spreadsheet with metrics and deliver it to the CISO's desk.
 - B. Send the CISO a weekly email report with those metrics entered manually.
 - C. Provide the CISO with logins to every individual security tool to find the data.
 - D. Build a real-time dashboard that pulls data from security tools to display these metrics.
18. A security engineer writes a Python script that pulls logs from a cloud provider every hour and pushes them into the SIEM. Which integration method is being used?
- A. Plug-in
 - B. Webhook
 - C. API
 - D. Manual integration
19. What does the MTTD metric measure in a cybersecurity incident response context?
- A. The total amount of uptime a system has before any incident occurs.
 - B. The average time from the start of an incident to the moment it is first identified by the security team.
 - C. The average number of security incidents detected in a given time period.
 - D. The average time it takes to fully remediate an incident after detection.

20. A new AI-based threat detection system sometimes classifies normal files as malware and even reports attack behaviors that never actually occurred. What risk does this illustrate?
- A. Data exposure
 - B. Model poisoning
 - C. Prompt injection
 - D. Hallucination

Answers to Review Questions

1. B. Developing a standardized playbook for malware incident response is best because it ensures all analysts follow the same procedures, which reduces confusion, increases efficiency, and produces consistent results. Allowing each analyst to use their own approach keeps the problem of inconsistency. Assigning all malware incidents to one person removes flexibility, overloads a single analyst, and creates a single point of failure. Delaying response until a manager advises wastes critical time and could allow the malware to spread or cause more damage.
2. D. AI models can learn from new data and adapt to novel attack patterns more quickly than static, signature-based systems, which helps detect threats that don't match existing rules. Forcing attackers to revert to old techniques is unrealistic since attackers continually evolve. Fixed rules would not provide the adaptability AI offers, so saying AI only operates that way is incorrect. Ignoring anything not in the initial database misrepresents how AI works, since its value lies in identifying anomalies and patterns beyond known signatures.
3. A. A documented, standardized incident response plan that defines roles and escalation points prevents wasted time because everyone already knows their responsibilities and when to involve higher authority, allowing the team to act quickly. Expecting team members to sort out roles during an incident leads to confusion and delays. Relying on management guidance for every decision slows down response and undermines autonomy. Simply adding more people does not solve the problem of unclear responsibilities and may actually increase confusion.
4. C. Correlating thousands of alerts from multiple tools to identify an attack pattern is best suited for AI because it involves rapidly processing and connecting large volumes of data, something machines do far faster than humans. Writing postmortem reports with cultural recommendations requires judgment and communication skills beyond AI's scope. Defining acceptable risk tolerance is a strategic business decision that must come from leadership. Writing customized firewall policies tailored to business needs demands understanding of context and priorities that AI alone cannot provide.
5. D. Automating the repetitive steps with a security orchestration and automation tool improves efficiency by freeing analysts to focus on tasks that require judgment while ensuring consistency and speed. Abandoning the standardized process removes the structure that already solved inconsistency issues. Assigning an intern to do repetitive work manually is

inefficient, error-prone, and not scalable. Having analysts continue those tasks wastes their expertise and reduces the overall effectiveness of the team.

6. D. This policy mitigates the risk of data exposure because pasting sensitive logs or confidential information into public AI tools could allow that data to leave the organization's control and be accessed by unauthorized parties. Malicious prompts refer to attempts to manipulate AI output, which is unrelated to protecting confidential data. Model poisoning involves corrupting the training data of an AI system, not user input leaks. Hallucinations describe AI producing false or made-up information, which does not involve the risk of exposing internal data.
7. C. Automatically isolating a workstation when it is confirmed as infected with malware is ideal for automation because it is a repetitive, well-defined action that must happen quickly to contain threats. Brainstorming new threat scenarios requires creativity and cannot be scripted. Interviewing an employee involves human interaction and judgment that automation cannot replicate. Developing a custom strategy for a zero-day attack requires expert analysis and adaptability, which automation cannot provide.
8. A. The AI Act governs the use of AI in the European Union, including its application to security-related work, by setting rules for transparency, accountability, and risk management. GDPR focuses on personal data protection, not specifically AI regulation. PCI DSS applies to securing payment card data and is not an AI framework. HIPAA regulates healthcare data in the United States, so it does not apply to AI use in the EU.
9. D. Alerts enriched with current threat context provide analysts the primary benefit because they can more effectively prioritize which threats are real and most urgent, improving response quality and speed. Automatically identifying and blocking threat actors across all systems is unrealistic because intelligence feeds may contain false positives and require validation. Dynamically rewriting firewall rules introduces risk of misconfiguration and disruption. Reducing the need for endpoint protection software is inaccurate since layered defenses are still necessary regardless of threat intelligence integration.
10. A. Using a webhook lets one system send an immediate HTTPS request to another to start the scan as soon as a new item arrives, which supports near real time action. Manual exports introduce delays and errors. An email forces humans to intervene and slows the process. A weekly PDF workflow reacts too slowly to critical issues.
11. A. Infrastructure as code is the best choice because it allows security settings like firewall rules and logging to be defined in templates so that every new virtual server automatically inherits them, ensuring consistency and compliance. A SIEM focuses on collecting and analyzing logs rather than enforcing configuration. A SOAR platform automates response to incidents but does not control how servers are initially built. Threat hunting involves proactively searching for adversary activity and does not address enforcing standard server configurations.
12. A. Validating AI output is the best practice here because the team reviews the tool's findings before acting, which helps prevent mistakes caused by false positives. Blindly trusting the model's judgment risks acting on incorrect alerts. Fine-tuning the model could improve

accuracy over time, but that is not what is happening in this situation. Removing humans from the loop would increase efficiency but also increase the chance of errors going unchecked.

13. B. Implementing automated alert tuning to enrich alert data and filter out known benign events is the best approach because it reduces false positives, decreases analyst fatigue, and ensures attention stays on real threats. Increasing sensitivity would create even more false positives and worsen the problem. Hiring more analysts only adds cost and does not solve the root issue of noisy alerts. Replacing all detection tools with a single vendor solution risks reducing coverage and does not guarantee fewer false positives.
14. D. Using infrastructure as code templates that call provider APIs during provisioning ensures every VM gets logging agents and baseline firewall rules automatically and consistently, making the process repeatable and integrated with cloud tooling. A shared spreadsheet only provides guidance and still depends on humans to follow steps correctly. Opening a ticket for IT Operations adds manual effort and delays. Allowing each engineer to create one-off scripts leads to inconsistent results and lacks centralized control.
15. A. Fewer false-positive alerts are the expected result because suppression rules and models that recognize safe behavior filter out noise, letting analysts focus on suspicious activity that matters. False negatives can still occur since no detection system is perfect. Increasing the number of alerts would contradict the purpose of suppression. Eliminating the need to maintain detection rules is unrealistic because ongoing tuning and updates are always required to adapt to evolving threats.
16. D. This describes data enrichment because the analyst gains additional context like WHOIS and reputation details attached to the IP address, helping make faster, better-informed decisions. Event correlation involves linking related alerts or logs across systems. Log analysis focuses on examining raw log data for insights. Tool orchestration refers to automating workflows between different security tools rather than simply adding contextual information.
17. D. Building a real-time dashboard that pulls data from security tools is best because it provides the CISO with immediate visibility into critical metrics like MTTD, MTTR, and alert volume without manual effort. Printing spreadsheets or sending weekly email reports introduces delays and requires constant manual updates. Giving the CISO logins to every tool is inefficient and forces them to navigate multiple interfaces instead of seeing consolidated insights in one place.
18. C. This is an API integration because the Python script queries the cloud provider programmatically to pull logs and then sends them to the SIEM. A plug-in would be a ready-made module installed directly in the tool. A webhook would push events automatically in real time rather than requiring scheduled pulls. Manual integration would require a person to move the logs themselves instead of an automated script.

- 19.** B. Mean time to detect (MTTD) measures the average time from when an incident begins until the security team identifies it, showing how quickly threats are spotted. Uptime before an incident is unrelated to detection and does not measure security responsiveness. Counting the number of incidents detected describes volume, not time. The time to fully remediate after detection is measured by mean time to respond, not MTTD.
- 20.** D. This illustrates AI hallucination because the system is generating incorrect or made-up results, such as labeling benign files as malware or inventing attack behaviors. Data exposure refers to leaking sensitive information, which is not happening here. Model poisoning involves tampering with training data to skew results, which is different from simple misclassification. Prompt injection applies to manipulating AI through crafted input, which is not the case in this scenario.

Answers to Lab Exercises

Playbook	Provides manual orchestration of incident response. For example, specific incidents and threats each have their own.
Runbook	Used by IT operations for reference for routine procedures administrators perform.
SOAR	Platform that offers opportunities to improve your organization's use of threat intelligence. By bringing information about emerging threats into this platform, you can enrich data about ongoing incidents and improve your ability to react to emerging cybersecurity situations.
IaC	With this, configuration and policy settings are stored in version-controlled code rather than being applied manually in consoles or GUIs. This approach allows organizations to deploy security controls in a consistent, repeatable manner and to roll out changes to thousands of systems with a single update.
Dashboard	Gives security analysts, managers, and executives a real-time view of the organization's security posture. These dashboards can pull data from SIEMs, SOAR systems, vulnerability scanners, and threat intelligence feeds, presenting them in a unified view.

APIs	Are a primary means of integrating diverse security tools. These are programmatic interfaces to services that allow you to interact with that service without using web-based interfaces.
Webhooks	Allow us to send a signal from one application to another using a web request.
Pug-ins	Provide an opportunity to increase integrations. These are small programs that run inside other programs, adding additional functionality.

