

DEFENSIVE STRATEGIES

Recommendations for
Security Leaders for 2026



Embrace a 'Breach Assume, Verify, Contain' Mindset

Embrace the “Breach Assume, Verify Contain” Mindset

With stealthy attacks designed to evade detection, it’s safer to assume your preventive tools will fail occasionally. Shift your strategy from purely threat prevention/detection to also include proactive breach detection. This means adopting an assume breach mentality: operate as though an attacker might already be inside, and focus on finding and containing them quickly. Key tactics:

Continuous Threat Hunting: Establish a dedicated threat hunting function (or leverage a managed service) that continuously looks for signs of compromise without relying solely on automated alerts. Hunters should sweep for anomalies in network traffic, unusual user behavior, odd new processes or scheduled tasks, etc. Essentially, look for the subtle footprints an AI-assisted intruder might leave. By hunting as if an attacker is already present, you greatly improve the chances of catching stealthy threats.

Out-of-Band Validation: Don’t trust a “clean” EDR dashboard at face value. Use secondary scanning tools and periodic audits to validate that systems are truly clean. For example, run offline malware scans or cloud-based scans on a random sample of endpoints each month; use memory forensics on critical servers to see if anything is hiding only in RAM; deploy network threat detection that might catch suspicious traffic the endpoint agent misses. In other words, trust but verify your primary defenses. This can catch scenarios where an attacker has blinded or bypassed your normal sensors.

Deception and Canary Objects: Consider deploying deception technology – e.g. honeypot accounts, fake file shares or credentials, canary tokens embedded in sensitive documents – that an attacker would likely interact with if they’ve gotten in. These traps act as high-fidelity tripwires for intrusions. An AI-evasive attacker might slip past normal logs, but if they unknowingly touch a decoy resource (say, a fake administrative share or a dummy database entry), you get an immediate alert. Such measures can dramatically shorten detection time for silent breaches.

By assuming a breach could be happening right now, your team stays vigilant and is more likely to catch the subtle, AI-assisted threats that automated systems miss. The goal isn’t to foster paranoia, but preparedness – being ready to react swiftly and effectively when something is amiss, even if your tools are telling you “all clear.”

Diversify Your Defensive Stack– Avoid the Monoculture



Diversify Your Defensive Stack

Avoid the Monoculture

Revisiting the monoculture issue: do not put all your security eggs in one vendor's basket. While vendor consolidation has benefits (integration, cost, simplicity), it also creates a single point of failure. CyberStash recommends:

Layered Defense with Differing Technologies: If you primarily use one suite (say, Microsoft or CrowdStrike), introduce at least one complementary layer that is independent. For example, if you have an XDR agent from Vendor A, consider adding a network-based anomaly detection system from Vendor B, or a cloud security monitoring tool from Vendor C. These heterogeneous tools will each have different blind spots – but it's unlikely an attacker can evade all simultaneously if they work differently.

Cross-Checking Telemetry: Use integrations or an XDR platform that can ingest data from multiple sources (not just the suite's own components). If your endpoint agent misses something, maybe your DNS monitoring or cloud logs catch an oddity. Ensure that visibility overlaps. In practice, this could mean retaining certain best-of-breed tools even if you adopted a platform – e.g., keep that specialized email threat protection running alongside Microsoft 365's filters, or maintain a third-party vulnerability scanner even if your platform claims to have one.

Supplier Risk Assessments: Many organisations assume their smaller suppliers represent the greatest supply-chain risk. In reality, the opposite is often true. Evidence consistently shows that the largest vendors carry the highest risk because they operate with the broadest attack surface and rely on hundreds—sometimes thousands—of their own downstream suppliers.

Take Salesforce as an example. You may consider Salesforce a trusted supplier from a security standpoint, but how can any organisation of that scale adequately manage the security of its vast and complex supply chain? The point is simple: the larger the vendor, the larger the supply-chain risk—not the smaller.

Treat every security vendor as an extension of your own environment—because if they are compromised, so are you. Regularly review the privileges your security tools have. For example, could a vendor-initiated update deploy code across your entire fleet? What external services or endpoints do your supply-chain vendor applications communicate with? If you want to avoid a SolarWinds-type supply-chain breach, tightly restrict the domains and IP addresses your application servers are permitted to access, and validate every update in a controlled environment before it reaches production.

Red-Team Your Environment with AI-Enhanced Attacks



Red-Team Your Environment with AI-Enhanced Attacks

Testing your own defenses has never been more important. Traditional red teaming (hiring ethical hackers to simulate attacks) should evolve to incorporate AI techniques. We recommend:

“Retrofit” Penetration Testing: In addition to normal pentests, specifically task your testers (or internal red team) to use known exploits and attack methods augmented with AI – just like real adversaries now do. For example, have them take a public exploit or phishing email and use an AI tool to obfuscate it or morph it into ten variants, then see if your security tools catch these. This will reveal if an older issue could still nail you because of a slight change.

Adversary Simulation Exercises: Use frameworks like MITRE ATT&CK but include emerging AI-assisted techniques in the scenarios. For instance, simulate a deepfake voice call to the helpdesk – do they have a protocol to verify identity beyond the call? Or simulate malware that rewrites itself after initial detection – can your SOC handle an alert that “disappears” or changes? Tabletop exercises can also cover these: walk through how you’d respond if your EDR went blind, or if widespread false negatives were suspected.

Purple Team with ML Tools: Purple teaming (collaboration between red and blue teams) can incorporate machine learning tools. Red teamers might use AI to generate myriad phishing lures; blue teamers can use AI to analyze those and improve detection rules. By experimenting internally, you both sharpen your team’s skills and pre-empt what attackers might throw at you. Some advanced organizations are even leveraging AI to generate benign “noise” traffic during tests to see if their detection can pick out the attack – mirroring how attackers hide in the noise.

The objective is to find your weaknesses before the bad guys do, especially those introduced by AI-enhanced tactics. This proactive approach not only uncovers technical gaps but also procedural ones (e.g., does your incident response playbook account for an attack that has no malware file to isolate?). After each exercise, feed the lessons back into improving controls and training your staff.

Strengthen Identity, and 'Out-of-Band' Controls



Strengthen Identity, Validation, and “Out-of-Band” Controls

When attackers can imitate trusted channels and identities (whether via deepfakes or stolen tokens), it’s crucial to have verifications that are hard for AI-driven attacks to circumvent. Key measures to implement:

Out-of-Band Verification for Critical Actions: For any transaction or request that can cause major impact (financial transfer, changing credentials, data dump, etc.), require a second verification channel. For example, if a CEO supposedly calls or emails to request a money transfer, have a policy that a secondary confirmation (like an SMS or a callback on a known number) is mandatory. This simple step has foiled real-life deepfake scams. Ensure staff are trained and empowered to enforce it, even if the requester seems senior – better to annoy a VIP with an extra call than to wire money to criminals.

Identity Proofing and Monitoring: Strengthen your authentication processes. Deploy phishing-resistant MFA (like FIDO2 security keys or certificate-based auth) which are harder for AI-driven phish to steal than simple OTP codes. Monitor for impossible travel or odd access patterns (even if credentials are correct, AI might slip by but a human logging from two countries in 1 hour is a red flag). Zero Trust principles should be embraced: continuously validate that a user or system should have access, not just at login but throughout a session. AI-based attackers excel at using valid credentials, so the defense has to shift to detecting anomalies in usage of those credentials.

AI-Evidence Checks: Develop ways to detect AI-generated content if possible. There are emerging tools to analyze audio for deepfake characteristics, or text for likely AI origin. While not foolproof, combining such tools into your processes can add friction for attackers. For instance, if your company is frequently targeted by spoofed emails, an AI content filter might tag an email as likely machine-written – which can then trigger extra manual review before acting on it. Likewise, consider requiring video call participants (especially in sensitive meetings) to turn on cameras and perhaps use agreed-upon gestures or codewords (since deepfakes often struggle with certain real-time interactions). These are new kinds of protocols that might become standard as awareness grows.

Monitor for Abuse of Unusual Tools and Languages

Monitor for Abuse of Unusual Tools and Languages

Given that attackers are experimenting with languages and tools less familiar to enterprise defenders (like the Zig, Go, Rust malware trend), organizations should update their monitoring to catch anomalies in the development environment and tool usage:

Process Monitoring for Dev Utilities: If your environment suddenly sees a compiler or interpreter running that's never been seen before (e.g., a zig.exe compiler appearing on a server, or a Python instance on a machine that typically doesn't develop software), that should generate an alert. Many SOCs focus on malware behavior but not on toolchain behavior. Yet, catching the presence of a rarely-used compiler might tip you off to an attacker trying to craft payloads on your systems.

Script and Macro Auditing: Tighten controls and logging around use of scripting engines (PowerShell, wscript, cscript) and Office macros. We know attackers will continue abusing these, with AI making it easier to create obfuscated scripts. Implement strict policies: for example, only signed company-approved PowerShell scripts can run; all Office macros are blocked unless explicitly allowed; command-line logging is enabled to capture one-liner commands that AI-generated malware might execute. It's about reducing the noise and having visibility when built-in tools are misused.

File and Memory Analysis Using AI: Fight fire with fire – deploy AI-driven analysis to detect AI-evasion tricks. Modern EDR/MDR services (like CyberStash's own Eclipse platform, for instance) are incorporating machine learning models that can spot, say, when code is hiding in an image or stuffed in a data section of an executable abnormally. These help catch those Zig Strike-like techniques. Ensure your tools are updated to look for things like anomalous section names in binaries, self-modifying code behavior, or programs that access AI APIs (why would a Word document need to call an AI service?). Such telemetry might indicate an AI-augmented attack in progress.

Additionally, keep an eye on threat intelligence for newly popular languages or frameworks in malware. If tomorrow attackers pivot to, say, a spike in Rust or a niche scripting language, be ready to add detection rules for executables with those traits. Anomalies are often a sign of either innovation or error – in security, both warrant attention.

Invest in AI for Defense – But Do So Responsibly



Invest in AI for Defense – But Do So Responsibly

To cope with AI-enhanced threats, organizations should leverage AI/ML in their defense – but with caution and clear objectives:

Augment Analysts, Don't Replace: Use AI to handle grunt work: log correlation, anomaly flagging, even automated responses for known patterns. Free your human analysts to focus on complex, creative investigations (which AI attackers find harder to evade). For example, an AI system might cluster thousands of alerts and say “these 5 machines show similar strange behavior”; a human can then examine that pattern holistically. This pairing can markedly improve detection and response times.

Maintain Human Oversight: Always have humans in the loop for critical decisions. If an AI says “this is malicious, quarantine it”, have a seasoned analyst review unless it's absolutely confidently known. This prevents an attacker from trivially tricking your AI into false positives that cause self-harm (imagine an attacker getting your AI to quarantine a vital system process by spoofing telemetry – it's a new attack vector). Regularly sanity-check your AI's outputs. If it misclassifies something, feed that lesson back in (update the model or rules).

Secure Your AI Models: Treat your detection models and threat intel ML systems as sensitive assets. Control access to training data, use versioning, and watch for concept drift (if your environment changes or attackers change tactics, your models need retraining). Also, beware of blind trust in AI-based threat intel from external sources – vet and corroborate with traditional methods. In 2026, we may see attempts to pollute community/shared AI threat models with bad data, so a zero-trust mentality even with your defensive AI is prudent.

Ultimately, AI for defense is indispensable to keep pace, but it's not a set-and-forget solution. It's more like a junior analyst that works 24/7 – incredibly useful, but it still needs supervision and training by senior security staff. CyberStash's perspective is that human expertise augmented by AI will beat either one alone. The winners will be orgs that find the optimal synergy, not those that think AI can magically solve security or those that ignore AI entirely.

Cybersecurity Culture and Training 2.0



Cybersecurity Culture and Training 2.0

Finally, address the human element with updated approaches:

Educate on AI Threats: Update security awareness training to include AI-related scenarios. Employees should see examples of deepfake videos, AI-written emails, and be taught skepticism of perfect communication. The classic advice “look for typos” is no longer enough – training should emphasize process (e.g. always verify requests, even if the message looks flawless) over content clues. Interactive drills can help, like sending simulated AI-generated phishing emails to see if users can spot any telltale signs or follow procedures.

Encourage a Reporting Culture: In an AI-rich threat landscape, things will slip through. It’s crucial that employees feel comfortable reporting anything suspicious, even if it’s just a gut feeling. Many successful breaches could be halted earlier if someone who noticed “that phone call felt off” or “this login at a weird time” speaks up. Reward vigilance and make it easy to report anomalies (quick IT/security hotline, anonymous if needed). Often humans observing context can catch what AI might miss – if they report it.

Limit Over-Sharing and Shadow AI: Train staff (especially developers and analysts) about the dangers of plugging company data into public AI tools. There have been incidents of proprietary code or data leaking via ChatGPT logs. Develop clear policies about AI usage: e.g., disallow entering sensitive info into external AI systems, use approved internal AI for such tasks if available, etc. This reduces the chance of an insider accidentally aiding attackers by leaking info that could train their models or reveal your secrets.

Also consider specialized training for the security team: ensure they stay up-to-date on AI developments in cybercrime, perhaps through courses or threat intel briefings. A well-informed team can adapt faster to new tactics as they emerge.

