

NEWSLETTER



THREAT INTELLIGENCE BY

ATHENA

TABLE OF CONTENTS

01	From the Editor's Desk	03
02	Latest Cyber Breaches	06
	a. Sinobi Ransomware Hits Indian IT Services Firm.....	06
	b. DragonForce Ransomware Hits Pantomath Investment Bank.....	07
	c. Axtria Data Breach: Analytics Platforms as Attack Vectors.....	09
03	Latest Cyber Scams	10
	a. Rs 100 Crore Digital Arrest Scam: SIM Box Network Exposed.....	10
	b. Rs 641 Crore Fraud: Chartered Accountants Turn Criminals.....	12
	c. Pan-India Cyber Fraud Crackdown: 27 Arrested in 11 States.....	14
	d. Southeast Asia Cyber-Slavery: A Structural Threat to India.....	15
04	Emerging Technology Trends	16
	a. AI Vulnerability Surge: First Deep-Exploit Mandate Emerges.....	16
	b. India's 3-Hour Deepfake Rule: A Global First in AI Law.....	18
	c. AI vs AI: Reshaping India's Cybersecurity Landscape.....	19
05	DIY Guide: Staying Safe in the Age of AI	21

From the Editor's Desk

Conditions beyond our control forced us to go on a brief hiatus. That brief period of disconnect from you, the reader, was sorely missed, and going forward, the chances of such instances recurring should, at best, be a thing of the past. Without wasting much time, let's get back to bringing the cyber world closer to you, so that you are better informed and thereby emboldened enough to stay at arm's length of digital harms.

While cyberattacks continued unabated across the globe, some of which are listed in this newsletter on the pages that follow, a more disturbing development occurred in the first quarter and spilt over into the second quarter of the year. This was primarily the Indian cyber space being under fire from hacktivist groups, owing allegiance to Palestine. While India was in no way directly involved in the war, it was affected, as were other countries dependent on energy imports. This impact was there for everyone to see, but there was a far more sinister operation at play, which not many were privy to, wherein hacktivists owing allegiance to the Palestinian cause, wreaked havoc primarily on websites of Indian organisations, including those of the Central Government and even certain state governments. More about that in the next issue, as action on that front is still underway. In the pages that follow, we have taken the time to keep you abreast of developments in the realm of cybersecurity so that you are well informed as you navigate the dangers of this hyper-connected world.

The first quarter of 2026 (Jan-March 2026) witnessed a significant transformation in the country's cyber threat environment. The first quarter of 2026 was marked by a significant increase in ransomware attacks, large-scale financial fraud schemes, and the growing use of artificial intelligence (AI) in both launching cyberattacks and bolstering cybersecurity defences.

Cyber risks are no longer isolated incidents affecting single organisations; today, they affect even the man on the street, and this trend is expected to intensify. If recent trends are any indication, nefarious elements active in the digital space are lately very prolific in exploiting digital interdependencies, human vulnerabilities, and emerging technologies. All this is done to amplify the scale and impact of cyberattacks. Apart from a lack of cybersecurity knowledge, another factor that has, in a way, strengthened the nefarious elements in their endeavour to exploit vulnerabilities is that defensive capabilities often lag behind due to operational complexity, slower response cycles, and reliance on traditional security models.

Some of the major cyber breaches that are covered in greater detail in this newsletter are the Sinobi Ransomware Attack on Impressico, the DragonForce Attack on Pantomath Group, and the Axtria Data Breach.

While cyber breaches in the first quarter of 2026 witnessed increased technical sophistication, cyber scams revealed a parallel evolution centred on psychological manipulation. These schemes exploited human emotions such as fear, urgency, trust, and greed, affecting individuals across various demographic groups. Some of the major scams that occurred during this period, and those covered in this newsletter, are the Digital Arrest Scam Syndicate (₹100 Crore fraud) and the ₹641 Crore Investment Fraud.

The first quarter of 2026 also witnessed a pan-India crackdown on cyber fraud. In March 2026, a coordinated law enforcement operation led to the arrest of 27 individuals across 11 states in connection with a nationwide cyber fraud network. The group was involved in various fraudulent activities, including investment scams, malware distribution, social media impersonation, credit card fraud, and visa-related scams. A concerning development covered in greater detail in this newsletter is the emergence of cybercrime compounds in Southeast Asia, notably in Cambodia, Myanmar, and Laos. These facilities have been identified as significant sources of cyber fraud targeting Indian citizens.

Some of the major developments in technology this quarter include AI continuing to reshape the cybersecurity landscape by enabling attackers to automate key aspects of cyber operations. Further, in February 2026, India introduced a comprehensive regulatory framework to address the risks associated with synthetic content and deepfakes. The framework defines synthetically generated information (SGI) and mandates the use of watermarks and metadata to identify such content. Another interesting development, especially in the Indian cybersecurity environment, is increasingly characterised by an “AI versus AI” dynamic. While organisations deploy AI-based tools for threat detection and response, attackers are leveraging similar technologies to enhance phishing campaigns, create deepfake impersonations, and exploit vulnerabilities.

We hope you enjoy sifting through the pages of this newsletter as much as we enjoyed putting it together, so you can read it at your leisure. Please do keep writing to us at pankaj.toppo@atheniantech.com, and we will look into each one of those suggestions so that we can improve this offering, which we bring only for you. Till then, happy reading ●

Patz

Pankaj Anup Toppo



Sinobi Ransomware Hits Indian IT Services Firm

On 28 January 2026, a ransomware group which goes by the alias of Sinobi, attacked Impressico Business Solutions, in one of the most sophisticated campaigns in recent years. The Sinobi group targeted the India-based IT services company's virtual machine infrastructure by attacking virtual machines and customer backup systems. This was a planned attack on the operational base of the managed services provider. The group claims to have stolen over 150GB of data, including contracts, financial information, and private customer records.

The attack vector followed a pattern that is becoming more common in advanced ransomware attacks on virtual infrastructure. By going after Hyper-V servers, the group encrypted not just single machines but whole virtual environments. The backup systems designed to protect them were also compromised. This resulted in a double blow, a calculated move to maximise operational damage, increasing pressure on the victims. This maximised operational damage is designed to force the victims into complying with the ransom demands by leaving them with no choice but to make a ransom payment or initiate complete system reconstruction.

What makes Sinobi different from normal ransomware actors is its deliberate targeting strategy. The Sinobi ransomware group has claimed to have around 50 victims globally in the opening weeks of 2026. The group specifically hunts for IT service providers and managed service providers (MSPs) as not just targets but multipliers. By breaching a single IT vendor's infrastructure, Sinobi gains potential access to hundreds of clients that the IT company provides services to.

The Sinobi attack did not happen in isolation. January 2026 has seen a broader and alarming surge in ransomware attacks across India, with cybersecurity researchers reporting that attacks rose more than 31 per cent above the previous nine-month average. Alarming, India accounted for approximately three per cent of global ransomware victims during the month. More than 38 active ransomware groups, including Qilin, LockBit 5, SafePay, Akira, DragonForce, and Sinobi, simultaneously targeted Indian organisations across professional services, manufacturing, IT, healthcare, and education sectors. For every organisation that relies on a managed service provider for its IT infrastructure, the message could not be clearer: the security posture of your vendor is, in effect, your own security posture.

DragonForce Ransomware Hits Pantomath Investment Bank



On 4 March 2026, Pantomath Group, one of India's biggest mid-market investment banks, appeared on the dark web leak site of the DragonForce ransomware group. The Pantomath group serves clients across mergers and acquisitions advisory, capital markets, structured finance, and asset management across more than 12 countries. What followed was not merely a data breach but a deliberate, surgically targeted exposure of extremely sensitive information. Cybersecurity researchers confirmed that the leaked data included internal corporate documents, financial spreadsheets, investment status reports, business development records, and internal email folders dating to March 2025.

DragonForce is not a new actor in the ransomware scene. Operating as a Ransomware-as-a-Service (RaaS) group, they have listed 363 victim organisations on their dark web leak site between their emergence in late 2023 and January 2026. Their activity has sharply increased from 2025 onwards, and what's worse, the group employs what the researchers call "double extortion". This involves simultaneously encrypting the victim's systems and threatening to publish the stolen data unless a ransom is paid by a defined deadline.

Latest Cyber Breaches

The targeting of Pantomath Group reflects a broader trend, with ransomware groups increasingly seeking out financial sector victims. This is not just for the ransom value but for the intelligence value of the data they hold. An investment bank sits at the centre of confidential commercial transactions. Its servers hold deal timelines, valuation methodologies, client identity, unreleased financial data, and strategic plans for companies across multiple industries. In the wrong hands, such information can be weaponised far beyond the immediate financial extortion.

DragonForce has demonstrated cross-platform deployment prowess across Windows, Linux, ESXi, and NAS environments, leaving many infrastructure stacks vulnerable. With its cartel-style expansion model and active recruitment on dark web forums, DragonForce is assessed by multiple threat intelligence providers as one of the most significant ransomware threats to Indian financial and professional services organisations in 2026.



Axtria Data Breach: Analytics Platforms as Attack Vectors

In January 2026, threat intelligence researchers identified that Axtria, a global provider of cloud software and data analytics solutions, had suffered a significant data breach. A threat actor published the company's proprietary data on a prominent hacking forum, claiming that internal business documentation, analytics assets, and client data had been exfiltrated and were available on the open dark web. The revelation placed a global analytics company squarely in the crosshairs of the cybercriminal underground.

Axtria's business model is built on collecting, processing, and analysing vast quantities of data on behalf of enterprise clients, particularly in the life sciences and pharmaceutical sectors. The company's platform handles commercial analytics, sales performance data, and market intelligence for some of the world's largest pharmaceutical companies.

The Axtria incident reflects a growing pattern in which data analytics firms are becoming attractive targets for sophisticated threat actors. Unlike a hospital or a bank, whose breach value to an attacker is immediately obvious, an analytics firm's breach value is less visible but potentially more strategically significant. This is because proprietary analytical models, commercial intelligence dashboards and client-specific insights contain years of intellectual investment.

Their exposure on a dark web forum constitutes not merely a privacy violation but a potential commercial catastrophe for the organisations whose data underpins those analytics. For organisations operating in the data-intensive technology sector, the Axtria case is a powerful reminder that the interest of the attacker is determined not by what you consider valuable, but what they can be enabled by your data.



Latest Cyber Scams

If the cyber breaches of the first quarter of 2026 were defined by the sophistication of their technical execution, the cyber scams that unfolded across the same period were defined by something even more unsettling and frightening. This is because these scams aimed at exploiting human emotion, psychological vulnerability, and institutional trust. January, February and March 2026 witnessed a plethora of fraud operations targeting Indians across every demographic.

These were from all walks of life, from retired professionals and small-business owners to senior executives. What unites these incidents, which are elaborated in detail below, is not their technical complexity but their psychological architecture. Each scam was carefully designed to trigger fear, greed, urgency, or affection, and each left its victims not just financially poorer but emotionally shattered.

Rs 100 Crore Digital Arrest Scam: SIM Box Network Exposed

On 10 January 2026, Delhi Police's elite Intelligence Fusion and Strategic Operations (IFSO) unit announced the dismantling of one of the largest and most technically sophisticated "digital arrest" operations ever uncovered in India. The syndicate, which had links spanning Taiwan, Cambodia, Pakistan, and Nepal, had defrauded victims across India of an estimated Rs 100 crore. The syndicate primarily operated through a network of 20,000 SIM cards, more than 5,000 compromised IMEI numbers, and 22 SIM box devices installed in cities including Mumbai, Mohali, and Coimbatore.



Latest Cyber Scams

The modus operandi was methodical and chilling. The syndicate's operatives, placed calls to victims across multiple states while posing as officers of the Anti-Terrorist Squad and other law enforcement agencies. Victims were told that arrest warrants had been issued in their names, connected to fabricated terror incidents including the Pahalgam attack and alleged Delhi blasts, and that their bank accounts were about to be frozen.

They were then subjected to continuous audio and video surveillance and instructed not to leave their homes or contact family members. This happened all the while being under the threat that any break in contact would be treated as an attempt to evade lawful custody. Under sustained psychological pressure, victims transferred money in multiple transfers to "verify their innocence."

Central to the syndicate's ability to evade detection was its deployment of SIM box technology. These devices contain dozens of SIM cards, which route international calls from Cambodia through local Indian numbers via 2G networks. These make them indistinguishable from domestic calls by circumventing standard telecom monitoring. The key technical operator, a Taiwanese national identified as I-Tsung Chen, 30, was arrested at Delhi's international airport on 21 December 2025.

He was apprehended after investigators ran a sting operation that convinced him the Indian network was still operational. Forensic analysis of seized devices was then conducted which included laptops, routers, CCTV cameras, foreign SIM cards, and Cambodia work visas.

The investigation, launched in October 2025 following a surge in complaints nationwide, led to seven arrests across Delhi, Mumbai, Mohali, and Coimbatore. The Supreme Court of India subsequently directed the CBI to prioritise investigations into digital arrest scams. This is the first time the country's highest court has directly intervened to address a specific category of cybercrime. In January 2026 alone, victims of digital arrest scams across India lost an estimated Rs 3,000 crore, according to India's Cyber Crime Coordination Centre (I4C) under the Ministry of Home Affairs.

Rs 641 Crore Fraud: Chartered Accountants Turn Criminals

In one of the most shocking cases of 2026, on 28 February 2026, the Enforcement Directorate arrested two chartered accountants, Ashok Kumar Sharma and Bhaskar Yadav, in connection with a Rs 641 crore cyber scam. The scam involved a financial fraud and money laundering operation that had victimised citizens across India. The case is significant not merely for its scale but for what it reveals about the professional infrastructure that enables high-value cyber scams. In this instance, it was not anonymous hackers but chartered accountants who operated the money laundering system that allowed Rs 641 crore in falsely obtained funds to flow undetected through India's financial system.

The scheme operated on a layered architecture that has become endemic to sophisticated investment fraud. Victims were approached via a digital platform that promoted attractive investment opportunities, part-time job offers, QR-code-based payment schemes, and phishing operations. Initial small payments were made to build trust and encourage larger investments, after which the victim's funds were frozen under various excuses. The proceeds were then routed through mule bank accounts managed through Telegram groups before being transferred to a UAE-based fintech platform called PYYPL. This effectively moved the money beyond the reach of Indian financial monitoring systems.



Latest Cyber Scams

A Chase That Spanned Three Courts

The investigation itself was a study in persistence. Enforcement Directorate searches were conducted at the residences of the accused in November 2024. Ashok Kumar Sharma allegedly fled his premises during the search, assaulting ED officials in the process, before both accused went into hiding and mounted a legal campaign to obtain bail. Their applications were rejected successively by the Special Court on 15 January 2025, by the Delhi High Court on 2 February 2026, and finally by the Supreme Court on 18 February 2026.

The ED has since arrested 10 individuals in the case, seized assets worth Rs 8.67 crore through two provisional attachment orders, and filed two prosecution complaints before the Special Court under the PMLA (Prevention of Money Laundering Act), 2002. This case serves as a reminder that some of the most dangerous cyber criminals wear suits and carry professional credentials.



Pan-India Cyber Fraud Crackdown: 27 Arrested in 11 States

In early March 2026, Delhi Police arrested 27 individuals across 11 states in a coordinated crackdown on a pan-India cyber fraud syndicate. While confirmed fraud amounts directly linked to the arrested individuals stood at over Rs 1.5 crore, forensic investigation of the suspects' bank accounts revealed transactions exceeding Rs 13.91 crore. Approximately 150 complaints were filed on the National Cyber Crime Reporting Portal, linked to their phone numbers, suggesting that the operation's reach was considerably larger.

The syndicate was engaged across a broad spectrum: investment fraud, APK file-based fraud, social media impersonation, credit card fraud, task-based fraud, and even visa-related fraud. Items recovered during the investigation painted a picture of a distributed, technology-enabled operation with both digital and physical infrastructure spread across the country. The syndicate was not concentrated in a single location but spread across 11 states, with different members handling different functions: SIM card procurement, bank account management, victim communication and fund extraction.

This geographic fragmentation is a deliberate strategy, designed to complicate investigation and prosecution by ensuring that no single jurisdiction holds the complete picture of any individual crime. The Delhi Police's ability to execute simultaneous arrests across state lines is a meaningful step forward in India's ability to pursue and dismantle such operations.



Southeast Asia Cyber-Slavery: A Structural Threat to India

Among all the cyber threats facing India in 2026, perhaps none is as structurally complex, geopolitically entangled, or humanely disturbing as the one emanating from the high-security scam compounds scattered across Cambodia, Myanmar, and Laos. According to data from the Ministry of Home Affairs, over 50 per cent of all cyber frauds targeting Indian citizens in 2025 and 2026 originated from these facilities.

The scale of the problem is staggering. I4C has identified at least 68 India-linked victims who collectively lost over Rs 12 crore to platforms affiliated with these compounds in recent months alone. Investigations by Gujarat Police revealed the inner workings of one such network where trafficked Indian nationals were held under intimidating conditions and forced to run fraud scripts targeting Indian citizens. Workers who refused faced physical coercion and confiscation of their travel documents. Those who complied became instruments of the very crimes destroying the financial lives of their fellow Indians.

What makes the Southeast Asia compound model so resistant to conventional law enforcement responses is its deliberate exploitation of jurisdictional boundaries. The criminals operate outside Indian law, in countries having varying degrees of cooperation with Indian authorities and INTERPOL. The financial flows are routed through layers of cryptocurrency transactions, shell companies, and overseas fintech platforms.

India's I4C has escalated cooperation with INTERPOL, ASEAN law enforcement bodies, and bilateral channels with Thailand, Vietnam, and the Philippines to address regional offshoots of these networks. The global crackdown on the Cambodia-based Prince Holding Group in 2025 demonstrated that coordinated international action can disrupt even the largest and best-resourced criminal networks. But dismantling individual networks is not enough. What is required is an intelligence-led, multi-national approach that treats cyber-enabled human trafficking and transnational fraud as connected and continuous threats.

AI Vulnerability Surge: First Deep-Exploit Mandate Emerges

Athenian Tech researchers say the security landscape is entering a new phase in which AI is no longer only assisting defenders, but also compressing the time from vulnerability discovery to exploitation into hours. This shift is reshaping enterprise risk by making machine-speed attacks more practical than human-paced patching and review processes.

At the centre of the warning is a new class of AI capability that can chain multiple weaknesses into a working exploit path with little human guidance. The researchers describe a system that can generate large numbers of exploits, uncover long-standing flaws in operating systems and file systems, and even behave in ways its creators did not explicitly intend.

The most important implication is that the traditional security timeline has collapsed. What once took days, weeks, or even longer can now happen in the span of a working day, which makes quarterly patch cycles and reactive vulnerability management increasingly inadequate.



Emerging Technology Trends

This creates a structural imbalance between defenders and attackers. Security teams are still often constrained by approval chains, ticket queues, and human review, while adversaries can use AI to discover, validate, and weaponize weaknesses at a far faster pace.

The researchers recommend a rapid executive response built around four priorities: expanding security capacity, mandating AI tooling across security operations, accelerating governance, and hardening environments through segmentation, phishing-resistant MFA, and stronger asset inventory. They also push for continuous vulnerability operations, AI-assisted code review, and deception controls that can respond at machine speed.

Their conclusion is stark: AI has changed the economics of offensive security faster than most organizations have adapted their defenses. The companies that move quickly will use AI to improve detection and remediation, while those that delay risk facing a world where the gap between disclosure and compromise is almost nonexistent.



India's 3-Hour Deepfake Rule: A Global First in AI Law

On 10 February 2026, the Ministry of Electronics and Information Technology notified the Information Technology (Intermediary Guidelines and Digital Media Ethics Code) Amendment Rules, 2026, effective from 20 February 2026. In doing so, India became the first jurisdiction in the world to enact a statutory framework that defines "synthetically generated information" (SGI), mandates permanent metadata and visible watermarks, and imposes takedown timelines for AI-created deepfakes. The amendment represents India's intervention into the governance of AI-generated content and its implications for the country's digital ecosystem are profound.

At the heart of the amendment is the introduction of a precise legal definition of SGI, i.e. any audio, visual, or audio-visual content created or altered algorithmically to appear real or indistinguishable from a natural person or real-world event. Deepfakes, AI-cloned audio, fabricated videos, manipulated images, impersonation material, and deceptive synthetic news content all fall within the scope. What was previously a vague, contested concept in Indian law is now a defined legal category with specific compliance obligations attached to it. Online intermediaries must now periodically notify users about the legal consequences of creating or sharing prohibited deepfakes, implement automated detection tools, and maintain audit logs that demonstrate compliance.

The most operationally significant provision of the amendment is the compressed takedown timeline. Government or court notices now require platforms to remove designated SGI content within three hours, a reduction from the previous 36-hour window. For the most sensitive categories like non-consensual deepfake nudity and content involving minors, platforms have just two hours from a user complaint to effect removal.



Missing either window carries a severe penalty: the platform immediately loses its safe harbour protection, making it as if it had created the content itself. The amendment also requires platforms to disclose the identity of deepfake creators to law enforcement where a crime has been committed, integrating the framework with the Bharatiya Nyaya Sanhita, 2023. The new rules are already creating a significant market for Indian AI-detection startups, which are being urgently engaged by social media platforms to manage compressed takedown windows.

AI vs AI: Reshaping India's Cybersecurity Landscape

A fundamental shift is underway in the nature of cyber conflict in India, one that cybersecurity professionals are increasingly describing as the arrival of an "AI vs AI" era. On one side, defenders: CERT-In now detects over 265 million security events annually, and Indian enterprises are deploying AI-powered behavioural analytics and anomaly detection. On the other side, attackers have embraced AI as an operational weapon by deploying AI-generated phishing emails and deepfake voice and video clones for social engineering.

The data is sobering. Studies indicate that one in six security breaches now involves AI in some form. AI-generated phishing is the most common, accounting for 37 per cent of AI-assisted incidents, followed by deepfake impersonation at 35 per cent. India's UPI infrastructure, processing over 15 billion transactions monthly, has emerged as a particularly high-value target for AI-enabled fraud. The Reserve Bank of India's Fintech Regulation and Ethics in AI committee has responded by mandating that banks and non-banking financial companies deploy AI-enabled monitoring, zero-trust architecture, and compulsory high-value fraud reporting.

Emerging Technology Trends

The broader regulatory and commercial response to the AI-driven threat landscape is driving a decisive shift away from traditional security toward zero-trust architecture, a framework in which no user, device, or network segment is trusted by default, and every access request must be continuously verified. Zero trust is rapidly transitioning from a new strategy to a regulatory baseline in India. Today's attackers are increasingly "logging in rather than hacking in", using stolen credentials and legitimate remote access tools to move through networks that perimeter defences cannot see.

Simultaneously, the emergence of "prompt injection" attacks has added an entirely new dimension to the security challenge. These attacks involve malicious inputs crafted to manipulate AI-powered customer service bots and financial automation tools into revealing sensitive data or authorising fraudulent transactions. Security audits must now include AI-specific red-teaming: hiring ethical hackers to attack an organisation's own AI systems. The lesson is clear: an AI system that is not security-audited is a glaring liability.



DIY Guide: Staying Safe in the Age of AI

We truly live in an age where artificial intelligence has become part of everyday life, and it has done so with remarkable speed. One moment it is helping us summarise research from a report, or generate an image, and the next it is being used as a virtual assistant, a search engine, a tutor, and even a work companion. With so much convenience packed into one place, it is only natural that people begin to share more than they should, often without pausing to think about what is actually being exposed.

Before getting into the details, it helps to understand a simple truth: modern AI tools are not private diary boxes. They are software systems that process what you type, what you upload, and sometimes what you allow them to remember. That means your habits matter as much as the technology itself. If you treat an AI assistant like a trusted colleague, you may already be giving away more than you intended.

In today's digital routine, people use AI for work, study, shopping, travel, coding, writing, and even personal advice. This makes it incredibly useful, but also easy to over-share. Many users paste long emails, screenshots, documents, resumes, medical notes, account details, and even personal frustrations into chat boxes without considering the downstream risk. The result is that a tool meant to assist can quietly become a place where your most sensitive information accumulates.

The first rule is simple: do not share Personally Identifiable Information (PII) unless it is absolutely necessary. Personally identifiable information includes your full date of birth, passport number, Aadhaar number, PAN number, bank account details, card numbers, CVV, passwords, OTPs, phone numbers, home address, exact workplace credentials, and private identity documents. Even seemingly harmless details, when combined, can create a complete profile of who you are and where you can be reached.

DIY Guide: Staying Safe on Modern AI

You should also be careful with anything that can be used for impersonation. That includes selfies with identity documents, utility bills, salary slips, employee ID cards, insurance papers, school records, and screenshots showing account balances or internal systems. If a task can be completed without exposing the original document, it is almost always better to redact or summarise first. A good habit is to ask yourself whether a stranger would be comfortable seeing the same information.

How People Get Exposed Through AI

One common mistake is treating AI tools like secure vaults. People often upload entire resumes, confidential work files, client data, legal drafts, or private chats in hopes of a faster answer. But the more context you provide, the more information you may be disclosing beyond the immediate task. Even when the tool is helpful, the input itself can become a privacy risk if it contains names, IDs, locations, or business-sensitive data.

Another risky habit is sharing emotional or highly personal content without filtering it first. People sometimes tell AI tools about their relationships, health conditions, financial stress, travel plans, or workplace disputes in a way they would never post publicly. The problem is not that the tool understands your feelings, but that your words may contain identifiers, timelines, or private facts that should stay off the record. In digital life, convenience can easily become overexposure.



DIY Guide: Staying Safe on Modern AI

Safe Habits When Using AI

The safest approach is to share the minimum amount of information needed to get the result you want. If you need help drafting a complaint, use placeholders instead of names, account numbers, or addresses. If you want help analysing a document, remove signatures, account identifiers, and any data that could reveal a person's identity or your organisation's confidential details. A small amount of editing before you paste can prevent a large privacy problem later.

You should also build the habit of using generic examples. Instead of saying "my bank account was debited on 14 March from branch X," say "a bank transaction was debited unexpectedly." Instead of uploading a full employee letter, summarise the issue in plain language. This gives the AI enough context to help while keeping the sensitive parts out of the conversation.

It is equally important to approach AI-generated advice with caution. AI can sound confident even when it is wrong, outdated, or incomplete. That means you should verify anything related to law, finance, health, security, or official processes before acting on it. In short, use AI for assistance, not blind trust.



DIY Guide: Staying Safe on Modern AI

A Few Rules To Follow

- Do not share passwords, OTPs, card details, CVV, UPI PINs, or authentication codes.
- Do not upload identity documents unless there is a verified, necessary reason.
- Do not paste confidential company data, source code, client records, or internal reports without permission.
- Do not reveal home address, travel plans, school details, or exact location unless required.
- Do not rely on AI memory features for private or sensitive information.
- Do not assume deleted chats are always instantly or fully removed everywhere.

A useful mindset is to treat AI like a public conversation in a controlled room, not a sealed confidential meeting. If something would make you uneasy on a screen visible to others, it probably should not be typed into an AI prompt either. That one habit alone can prevent most accidental disclosures.

The real goal is not to fear AI, but to use it with judgment. Modern AI can be powerful, efficient, and genuinely helpful, but it rewards users who are careful, precise, and selective about what they reveal. If you keep your prompts clean, your personal information minimal, and your expectations realistic, you can enjoy the benefits of AI without turning convenience into exposure.





Athenian Tech

Contact Us

 www.atheniantech.com

 <https://www.linkedin.com/company/athenian-tech/>

