

REGULATORY CHALLENGES TO INVESTIGATING AI-DRIVEN CYBERCRIMES

Naeem AllahRakha¹

Tashkent State University of Law, Uzbekistan

ABSTRACT

Purpose: The exponential growth of artificial intelligence (AI) technologies has revolutionized various sectors, including communication, healthcare, finance, and law enforcement. However, this same technology has also been co-opted by cybercriminals, giving rise to new forms of digital threats that are more autonomous, adaptive, and difficult to trace. AI-driven cybercrimes such as deepfake manipulation, automated phishing, AI-enhanced ransomware, and botnet attacks challenge conventional models of crime investigation, legal responsibility, and digital evidence admissibility. The aim of this research is to critically examine the regulatory challenges that hinder the effective investigation and prosecution of AI-driven cybercrimes.

The problem statement is that existing legal systems are not sufficiently equipped to regulate or respond to the unique challenges posed by AI technologies in cybercrime. Current laws often lack precise definitions for AI tools used in criminal activity, clear rules for assigning liability when autonomous systems are involved, and consistent procedures for handling digital evidence generated or manipulated by AI. The objective of the study is to:

1. Identify the specific features of AI-driven cybercrimes that complicate legal investigations,
2. Assess the adequacy of current international and domestic legal instruments, and
3. Propose policy and legislative recommendations for more effective regulation and enforcement.

The research question is: How can legal frameworks be adapted to effectively regulate and investigate cybercrimes involving artificial intelligence?

The significance of this research lies in the fact that, as the global community increasingly relies on AI-powered systems, the risks of unregulated or malicious AI applications in cyberspace grow exponentially. This research seeks to contribute to international cyber law discourse in building more resilient legal infrastructures for the digital age.

Design/Methods/Approach: This research adopts a doctrinal legal methodology supported by comparative legal analysis and qualitative document review. Primary and secondary legal sources will be gathered exclusively from publicly accessible official web portals of international organizations, government bodies, and legal institutions. To ensure the reliability and academic integrity of the data, the CRAAP Test (Currency, Relevance, Authority, Accuracy, and Purpose) will be applied to evaluate the quality of each source. Only open-source materials available in the public domain will be used, ensuring transparency and adherence to ethical research practices.

Findings : The study reveals significant regulatory challenges in investigating AI-driven cybercrimes through recent high-profile cases involving artificial intelligence and digital harm. First, there is a major lack of legal clarity regarding liability for the unauthorized use of data by AI systems. This is evident in the consolidated US copyright lawsuits against OpenAI and Microsoft, where twelve cases brought by prominent authors (such as John Grisham, Jonathan Franzen, and Ta-Nehisi Coates) and

¹ Chaudharynaeem133@gmail.com



media outlets including the New York Times argue that their copyrighted content was used without consent or compensation to train large language models (LLMs) like ChatGPT and Microsoft Copilot. These lawsuits highlight the challenge of establishing responsibility when AI systems are trained on vast datasets, often scraped indiscriminately from the internet. The fact that many of these cases were transferred to a single district court for consistency underscores the systemic uncertainty in current regulatory frameworks about how to handle AI-generated infringement and associated evidence.

Second, the manipulative power of deepfakes and generative AI technologies poses a growing threat to personal dignity, reputational rights, and public trust. The fake deepfake video involving Scarlett Johansson, in which she was depicted along with other Jewish celebrities in a fabricated message targeting Kanye West, underscores the dangers of AI-generated disinformation. This case reveals regulatory gaps in safeguarding individuals against AI-enabled identity theft, defamation, and synthetic media abuses. It also raises serious concerns about the admissibility and verification of AI-altered content as legal evidence and the absence of robust deterrents against such misuse.

Originality/Value: This paper advances the frontier of cyber law by addressing the underexplored intersection of artificial intelligence and cybercrime, offering a timely and original analysis of how autonomous systems are reshaping the nature of digital threats and legal responses. Its core value lies in critically identifying the gaps in current regulatory frameworks that struggle to keep pace with the rapid evolution of AI technologies, particularly in the context of evidence, attribution, and jurisdictional complexity.

The study's originality is further illustrated through real-world cases, including the unprecedented incident in April 2025 when a New York appellate court was confronted with an AI-generated avatar attempting to deliver oral arguments on behalf of a self-represented litigant. The judges' swift intervention and refusal to allow the avatar to continue highlighted a deep unease within the judiciary regarding the lack of procedural safeguards, authenticity, and legal accountability in AI-mediated legal interactions. This research proposes targeted reforms to ensure that AI-generated content is treated with appropriate evidentiary scrutiny, that liability regimes are adjusted to include algorithmic actors, and that transnational cooperation is bolstered through harmonized legal standards.

Keywords: AI Cybercrimes, Digital Evidence, Cyber Law, International Frameworks.

About the author

Mr. **Naeem AllahRakha** is a professor and founding member of the department of Cyber Law, Tashkent State University of Law. He is a professional lawyer and member of the prestigious Bar Councils. He is a member of the UN/CEFACT working group and serves as visiting faculty at several international universities.

INTRODUCTION

A single deepfake video can spread across the internet in minutes, destroying reputations and confusing millions of people before the truth is revealed (Pesetski, 2020). Imagine a world where a fake video can ruin someone's career, or a smart computer program can steal millions without a human hand guiding it. This is not science fiction; it is happening now through AI. This shows how AI, while powerful and useful, can also be turned into a dangerous tool for cybercrime (Rahman-Jones, 2025). Criminals now use AI to create fake identities, launch smart phishing attacks, and design ransomware that learns and adapts (Burton et al., 2025). The laws we have today were not built for crimes



committed by machines that can act on their own (Nerantzi & Sartor, 2024). Understanding how to regulate and investigate these threats is not just interesting, it is essential for protecting trust, security, and justice in the digital age.

Cybercrime has been a problem since the early days of the internet, but the rise of Artificial Intelligence (AI) has made it more complex and dangerous (Krishna, 2024). In the past, cybercrimes were mostly carried out by humans using simple tools like viruses, spam, or malware. Today, AI makes it possible to create smarter attacks that are faster, harder to detect, and able to adapt to defenses (Guembe et al., 2022). Examples include deepfake videos, automated phishing, and AI-powered ransomware. However, there is still little research on crimes that are directly driven by AI. Existing laws do not clearly explain who is responsible when AI systems cause harm, and they do not provide strong rules for handling AI-generated evidence.

We already know that AI is being used by criminals to create new forms of cybercrime that are faster, smarter, and more damaging than before. What makes the problem worse is that AI can act without direct human control, which makes it difficult to know who should be held responsible (Königs, 2022). Current laws were designed for crimes committed by humans, not machines, and this creates confusion in investigations and court cases (Mahardhika et al., 2023). For example, there are no clear rules on how to handle digital evidence that may have been changed by AI, or how to assign liability when AI tools are misused. We still need to understand how legal systems can adapt to these new challenges.

The debate on AI-driven cybercrime keeps returning to three pressure points: evidence, attribution, and coordination, yet these are the very places where law and policy fall silent. On evidence, proposals to adapt authentication rules for synthetic media and to treat AI outputs as a special class of electronic evidence show promise yet remain unsettled (Pantanowitz et al., 2024). Digital-forensics work documents weak standards for handling and validating manipulated data, and even decision-making biases during triage, suggesting courts risk admitting or excluding the wrong evidence (AllahRakha, 2025). On threat characterization, technical reviews show adversaries leveraging machine learning at scale for zero-day exploitation, adaptive malware, and federated tactics. EU landscape reports also confirm ransomware and DDoS dominance across sectors. However, they stop short of linking these findings to enforceable legal tests, leaving a gap between technical evidence and legal frameworks (Mohsen-dokht et al., 2024).

Botnet surveys highlight AI-driven command-and-control and detection arms races, underscoring the need for standards on provenance, audit trails, and model transparency (Hyslip & Pittman, 2015). On liability, doctrinal analyses argue current criminal law lacks clear theories for assigning fault when autonomous or semi-autonomous systems generate harm, and they diverge on whether to ground culpability in designers, deployers, or a new legal subject (Gaeta, 2024). Internationally, cooperation initiatives emphasize policy alignment against ransomware but provide limited procedural guidance for cross-border AI evidence and shared model-risk audits (Härmand, 2023). Most literature points out the challenges, yet stops short of offering workable tools such as chain-of-custody standards, reliability checks, or jurisdictional rules for AI-driven acts.

Many scholars agree that AI-driven cybercrime is a fast-growing threat, yet critical gaps remain as criminals exploit AI to outpace weak laws, unreliable forensics, and unprepared justice systems. They also discuss liability and the need for stronger cooperation across borders. These are strengths because they give a good overview of the risks. However, most studies stop at describing problems and do not offer practical legal solutions for courts or investigators. There is little work on how to adapt rules of evidence when AI content may be fake or manipulated. Few studies explain how to assign liability fairly when AI systems act without direct human control. Even less research connects technical advances in AI with concrete legal reforms. The objective of the study is to:



- Identify the specific features of AI-driven cybercrimes that complicate legal investigations,
- Assess the adequacy of current international and domestic legal instruments, and
- Propose policy and legislative recommendations for more effective regulation and enforcement.

The research question is: *How can legal frameworks be adapted to effectively regulate and investigate cybercrimes involving AI?*

The rapid rise of AI in cybercrime has outpaced the ability of laws and regulations to respond, creating urgent gaps in justice and security. AI-driven crimes raise new questions about evidence, liability, and international cooperation. If these questions are not answered, courts may struggle to deliver justice, and criminals may take advantage of the legal gaps. The importance of this research is that it provides a clear analysis of where the weaknesses are and suggests ways to improve legal systems. Its impact on the research field is to expand cyber law by adding specific focus on AI, which has not been studied in depth. It contributes new knowledge on how to handle AI-generated evidence and assign responsibility in cases involving autonomous systems.

METHODOLOGY

This study uses a qualitative research design because it focuses on laws, policies, and strategies rather than numbers or statistics. The object of the study is a broad set of cybercrime regulations and policies, while the sample specifically focuses on rules and responses related to AI-driven cybercrimes. Data was collected from Google Scholar using keywords such as AI cybercrimes, digital evidence, cyber law, and international frameworks, and from official websites to ensure authenticity. Only peer-reviewed journal articles and official legal documents were included. The CRAAP test was applied to check validity and reliability, focusing on sources that are current, relevant, accurate, and credible. Data was analyzed through doctrinal and document analysis methods. All information comes from the public domain, with full citation to original authors to ensure integrity. The study is limited because laws and technology keep changing, which may affect the future value of findings. It assumes that sources are reliable and represent the true state of knowledge.

RESULTS

This study looks at the problem of AI-driven cybercrimes and how legal systems can deal with them. The main question is how laws can be adapted to regulate and investigate crimes involving AI. The research explored issues such as evidence, liability, and international cooperation. It focused on cases where AI created new challenges for courts, investigators, and lawmakers. The aim was to find where current rules are weak and to suggest possible reforms. This is important because AI crimes are growing, and traditional laws are not enough.

The research found that the biggest challenge in AI-driven cybercrime is the lack of legal clarity (AllahRakha, 2024). Current laws do not say clearly who is responsible when AI is misused or when it acts in ways not directly controlled by humans (Christen et al., 2023). Courts face serious problems with digital evidence that may be altered or generated by AI, because it is difficult to prove what is real and what is fake (Runyon, 2025). International rules are also weak, as countries use different systems and often cannot cooperate fast enough to stop global cybercrime (Wang, 2024). This means that criminals can exploit AI technology without facing clear punishment, while victims may not get justice.



One interesting finding pertains to how quickly criminals adopt AI tools for their attacks (André Maio, 2025). For example, phishing emails powered by AI are far more convincing than older ones, making it harder for people to tell the difference. Another example involves deepfakes, which are no longer just entertainment but powerful tools for fraud and defamation (Lundberg & Mozelius, 2025). The research also shows that even courts and legal professionals are experimenting with AI, such as the unusual case of an AI avatar appearing in court. This shows that AI is not only a problem in cybercrime but also in the justice system itself (Daniel, 2025). Another surprising point is that while governments talk about the use of AI, most have not updated their laws to match (Miranda, 2025). This gap between awareness and action creates serious risks. These findings highlight that AI is changing not only crime but also how justice works, and that reforms cannot be delayed.

The study shows four main findings. First, AI-driven cybercrimes are more complex and harder to investigate than traditional cybercrimes. Second, current laws do not clearly explain how to handle AI-related evidence or assign liability when autonomous systems cause harm. Third, international cooperation is weak, and this makes it harder to deal with crimes that cross borders. Fourth, there is a gap between the speed of AI development and the slow pace of legal reform. Together, these findings show that legal systems are not yet ready for the new challenges created by AI. The results make it clear that there is an urgent need for reforms that focus on evidence, liability, and cooperation. Without these, AI will continue to give criminals an advantage over justice systems worldwide.

An unexpected finding was the strong reaction of judges to AI inside the courtroom (Fine et al., 2023). The case of an AI avatar trying to present arguments in a New York court revealed deep fears in the legal system itself. This was not expected, because most research focuses on criminals using AI, not legal professionals confronting AI directly. Another unexpected finding is how much copyright disputes, like the lawsuits against AI companies for data use, connect to cybercrime debates (Montgomery, 2025). While they are not crimes in the traditional sense, they show how unclear laws are when AI uses data without permission. This link between intellectual property law and cybercrime was not obvious before. It shows that the impact of AI is wider than expected and touches many areas of law.

The research question asked how legal frameworks can be adapted to regulate and investigate AI-driven cybercrime. The results show that laws must change in three main ways. First, liability rules must be updated so that responsibility is clear when AI systems cause harm. Second, courts need stronger rules for handling AI-generated or manipulated evidence, including standards for verifying digital content. Third, international cooperation must be improved, since AI-driven crimes often cross borders and no single country can handle them alone.

DISCUSSION

When AI Becomes the Criminal Tool

Once deployed, AI tools can act independently. This autonomy makes them harder to trace back to the criminal (Caldwell et al., 2020). Evidence from phishing cases shows that AI-generated emails are nearly indistinguishable from genuine communication, reducing the ability of victims to detect fraud. While the strength of this evidence lies in real-world cases, its weakness is that many studies remain descriptive rather than quantitative. This means we lack strong statistical measures of impact. Still, the consistency across regions strengthens the evidence that AI has changed the game. Unlike past cybercrimes, these attacks are not just human-driven. They involve systems that partially think and act on their own. This raises deep questions about responsibility in law.



AI crimes are unique because they combine autonomy, speed, adaptability, and deception. Autonomy allows attacks to run without constant human input. Speed means AI systems can attempt thousands of intrusions in seconds, overwhelming defenses. Adaptability allows malware to switch tactics instantly if blocked. Deception shows up in deepfakes and AI phishing, designed to fool even experts (Park et al., 2024). Research on ransomware and deepfakes confirms these risks, though controlled lab studies often fail to capture the messy reality of cybercrime. While weak security or human error may also explain some breaches, repeated findings suggest AI itself is a new risk factor that amplifies criminal power (Završnik, 2020).

Investigations are especially difficult because AI blurs intent and liability. If an AI adapts mid-attack, was it by design, or did the system ‘decide’ on its own? This uncertainty makes it harder to assign blame (Fernandez-Basso et al., 2025). Forensic work also struggles to separate genuine evidence from AI-manipulated content, reducing trust in digital proof. In Minnesota, a judge struck expert testimony relying on AI-generated citations in a deepfake-related case (Thomas, 2025). This demonstrated how even legal experts can be misled by AI hallucinations; similarly, in Maryland, a high-profile case involved an AI-generated deepfake audio of a school principal, complicating the chain of custody and evidentiary reliability (Lake, 2024). Across the Atlantic, Europol’s 2025 Serious and Organized Crime Threat Assessment warns that AI-enhanced organized crime featuring voice cloning, deepfakes, and synthetic fraud has made digital evidence even more volatile and deceptive throughout the EU (Aleksandrowicz, 2025).

Court systems themselves are unprepared: experts warn that mounting sophistication in AI content makes verifying authenticity harder, as courts lack both sufficient forensic tools and trained personnel. Most public discussion centers on these headline-making episodes, while smaller-scale crimes like manipulated evidence in custody disputes or low-budget fraud go unreported. Earlier cybercrime research centered on human intent; today’s reality is more complex, where machines are active participants in the crime (Farooq & De Vreese, 2025). Recent legal disputes make this clear. Lawsuits over AI trained on copyrighted data show harm caused directly by systems, not just by their users.

Why Today’s Laws Fall Short

Domestic and international laws are not built for AI-driven crime. Most frameworks were written long before AI gained autonomy. As a result, they struggle to handle crimes committed or supported by intelligent systems. Recent copyright litigation such as *Bartz v. Anthropic*, where Judge William Alsup ruled that using purchased books for AI training may be fair use, but storing millions of pirated copies is not exposed a clear legal gap (Sherman & Hooker, 2025). Judges struggle to apply decades-old statutes to novel harms emerging from AI. International agreements, such as the Budapest Convention, support cross-border cooperation. Yet they offer little guidance on AI-generated evidence or algorithm-driven liability. The evidence is credible because it comes from real cases and treaties, but weak because it reflects early debates rather than lasting solutions (Carpenter, 2025).

Current laws focus on human action, intent, and harm. This model breaks down when AI generates phishing emails or deepfakes without direct human input. Judges often search for the “human actor behind the machine”, but this approach leaves gaps. Many AI-driven crimes involve minimal or indirect involvement by people. Case reviews show a consistent pattern: courts hesitate to treat AI as anything more than a tool, even when systems act with independence. This reflects a global legal bias toward human accountability, but it leaves crimes committed through AI in a gray zone (Daniele, 2024).

Another problem is the lack of clear definitions. What does “autonomous decision-making” mean in law? Without precise language, courts treat similar cases differently, which erodes trust in justice.



Ongoing lawsuits against AI companies show deep disagreement: is AI output original work, manipulated data, or stolen property? These conflicting interpretations make it harder to assign liability to designers, users, or corporations. Evidence is strong because it comes from real disputes, though most involve civil law, so lessons may not transfer cleanly to criminal justice (Nastoska et al., 2025).

Global inconsistency makes things worse. Some countries are quick to update their laws, while others lag behind. This creates loopholes that criminals exploit. A deepfake scam may be illegal in one country but ignored in another, offering safe havens for offenders (Sandoval et al., 2024). Reports from ENISA and the Council of Europe confirm that cross-border cooperation often collapses when data moves between jurisdictions with mismatched standards. The evidence here is strong, but biased, since many countries underreport AI-related crimes.

U.S. lawsuits against OpenAI and Microsoft test whether training AI on massive datasets is theft or fair use (Chen, 2025). Courts remain divided, revealing a major gap in ownership and liability rules. Deepfake misuse offers another example: fabricated videos used for fraud or reputational harm often outpace the legal system's ability to verify or punish. These cases are compelling because they involve high-profile courts and global attention. Still, most rulings are pending, so long-term outcomes remain unclear. Some argue the issue is not weak laws but unsettled ethical norms. Either way, today's legal systems are not ready for the speed and complexity of AI-driven crime.

Truth and Trust in AI Evidence

Courts today face serious problems when dealing with AI-generated or altered evidence. Deepfakes, synthetic texts, and altered audio can look real, leaving judges unsure what to trust (Young, 2025). Many courts lack clear rules on how to admit or reject such material, which leads to inconsistent rulings. This uncertainty weakens justice, as outcomes depend more on discretion than on reliable standards. Digital forensics has not kept pace with AI's ability to manipulate content, raising the risk of both error and abuse (Analiza V Muñoz, 2025). Some courts may even reject genuine evidence out of caution, while others may admit fakes. Either way, truth suffers.

Authentication and verification are more difficult in the AI era because traditional forensic methods were designed for older forms of digital data, such as metadata, file integrity, or IP tracing (Malik et al., 2024). They often fail against AI-generated content, which leaves no conventional fingerprints. Investigators struggle to prove origin or authenticity. Detection tools exist, but they are limited, biased, and easy to bypass. Watermarking of AI output is still rare (Dathathri et al., 2024). Optimists argue that technology will eventually solve this, but for now the gap is critical. Without new forensic standards, courts risk losing the ability to trust digital evidence at all.

This creates two dangerous risks: admitting false evidence or rejecting true evidence. If manipulated material enters court, the innocent may be punished while the guilty escape. If authentic evidence is dismissed as fake, victims are denied justice. Current studies confirm how hard it is to separate genuine from fabricated AI content, leaving judges to rely on personal judgment. Compared to older digital forensics, this problem is far more serious, AI makes fakes almost indistinguishable from reality (Steven, 2021). When truth is uncertain, public trust erodes. Victims may hesitate to report crimes, while citizens lose confidence in legal institutions (López-Borrull & Lopezosa, 2025).

The risk is not only technical but moral. Justice depends on trust as much as law. Unlike earlier cyber challenges, AI manipulation is sharper and more dangerous, because it can strike at reputations, elections, even national security. While courts may eventually adapt, the rapid pace of AI leaves little time to catch up. Possible solutions exist but remain underdeveloped. These include mandatory watermarking of AI content, stricter chain-of-custody rules, advanced forensic tools, and expert panels to assess



AI evidence before trials. Such reforms would limit judicial discretion and reduce bias. Yet challenges remain: tools are still imperfect, and questions of cost and responsibility whether governments, courts, or tech companies should lead are unresolved. The path forward is clear: if courts want to remain credible, they must quickly adapt laws and tools to protect truth in the AI era.

Building Smarter Rules for the Future

The hardest question is who should bear responsibility when AI causes harm. AI tools complicate this because they act with autonomy. Some suggest liability should fall on developers, owners, or users. Others propose shared models. Evidence from deepfake scams and AI ransomware shows that victims often face injustice when liability is vague. The strength of this evidence is that it reflects real harm. Its weakness is that laws differ widely, so fairness cannot be applied equally across borders. Alternative explanations, such as weak security or careless users, exist but they do not fully explain AI's independent role. Compared to older cybercrime, AI introduces new ethical and legal dilemmas. Effective liability models must balance fairness, cost, and deterrence to protect both victims and society (Custers et al., 2025).

Evidence is the backbone of justice, yet AI undermines it. Deepfakes and synthetic content can fool even experts. Courts risk admitting false material or dismissing truth, both of which erode justice. Current rules assume human-created evidence and are not designed for AI. Recent fraud cases involving deepfakes confirm the danger, though much research remains technical rather than legal. Forensic tools exist but often fail against advanced manipulation. Unlike earlier digital evidence challenges, AI makes the problem more complex. New rules for authentication, verification, and presentation of AI evidence are essential to restore trust in courts (Popa et al., 2025).

AI crimes are global by nature. Criminals operate across borders, but legal systems remain fragmented (Fidler, 2025). This mismatch slows down investigations and leaves many victims without help. International ransomware attacks reveal the problem clearly: states fail to act together, and criminals escape (Meurs et al., 2024). Most data come from Western contexts, but the threat is worldwide. Experts argue bilateral agreements are enough, yet evidence shows they move too slowly for AI-enabled attacks. Global standards on liability, evidence, and cooperation are no longer optional, they are urgent. Without them, criminals will continue to exploit borders while victims remain unprotected.

Even inside the courts, readiness is low. Judges and lawyers need training to understand AI technology, evaluate evidence, and apply rules fairly (L. Kimbrough, 2025). Some of the problem comes from limited funding or outdated legal education, but the speed of AI innovation is the biggest factor. Earlier studies called for digital literacy; today the need is for AI literacy. Training, resources, and expert support will reduce errors, bias, and delays. Without this investment, justice systems risk falling behind. With it, they can become strong enough to face the digital future.

CONCLUSION

AI has become a powerful force in shaping the future of crime and justice. It is not only a tool for progress but also a weapon in the hands of criminals. AI-driven cybercrimes such as deepfakes, adaptive phishing, and AI-powered ransomware show how quickly technology can outpace the law. The research shows that the greatest risks come from unclear legal responsibility, weak rules for AI-generated evidence, and poor international cooperation. These challenges are linked because weak liability



systems make it harder to punish offenders, unreliable evidence makes it harder to prove guilt, and poor cooperation lets criminals cross borders freely.

The central claim of this study is that legal frameworks must evolve to address crimes powered by AI. Laws designed for human actions cannot handle machines that act independently or manipulate evidence. Stronger liability rules, better evidentiary standards, and harmonized international policies are essential to protect society from these new threats. The discussion of results shows that reform is possible and necessary. Smarter liability models would make responsibility clear. New evidentiary rules would strengthen trust in courts. International frameworks would close the gaps across borders. Training and resources for legal professionals would prepare justice systems for the challenges ahead. These measures work together to create a safer and more reliable system.

There are voices that insist current laws and corporate self-regulation are sufficient to manage AI crimes, but this belief ignores the scale and speed of the threat. However, evidence shows that AI develops too fast for outdated rules, and profit-driven self-regulation cannot protect the public from harm. Readers should see that adapting laws is not only reasonable but urgent. This study contributes to the growing field of cyber law by connecting AI-driven threats with practical reforms. The insights apply to governments, courts, businesses, and individuals. The findings suggest that protecting evidence, clarifying liability, and promoting cooperation are the key to a fair digital future.

There is still more work to be done. Future research should explore how liability can be shared between developers, users, and owners. More studies are needed on methods for verifying AI-generated content in court. International dialogue must continue to build global standards. Practical next steps include pilot programs for AI evidence protocols and training courses for judges and investigators. The same urgency that opened this study remains true at its end. If a single deepfake can damage a life in minutes, only strong and forward-looking legal systems can protect dignity, trust, and justice in the age of AI.

REFERENCES

- Aleksandrowicz, M. (2025, March 19). *Europol warns of AI-driven crime threats*. <https://www.reuters.com/world/europe/europol-warns-ai-driven-crime-threats-2025-03-18/>
- AllahRakha, N. (2024). Legal Frameworks for AI-Driven Cybercrime Prevention. *Uzbek Journal of Law and Digital Policy*, 2(6), 1–24. <https://doi.org/10.59022/ujldp.253>
- AllahRakha, N. (2025). AI and Corruption: Legal Liability in Algorithmic Decision-Making. *Access to Justice in Eastern Europe*, 8(3), 303–326. <https://doi.org/10.33327/AJEE-18-8.3-a000120>
- Analiza V Muñoz. (2025). *AI in the Crosshairs: Advancing Cybersecurity and Digital Forensics in the Era of Intelligent Threats*. <https://doi.org/10.13140/RG.2.2.34418.77769>
- André Maio. (2025). *Artificial Intelligence and Crime: The Dual Role of AI in Criminal Activity and Crime Prevention*. Zenodo. <https://doi.org/10.5281/ZENODO.16945903>
- Burton, J., Janjeva, A., Moseley, S., & Alice. (2025). *AI and Serious Online Crime* (CETaS Research Reports). Centre for Emerging Technology and Security. <https://cetas.turing.ac.uk/publications/ai-and-serious-online-crime>
- Caldwell, M., Andrews, J. T. A., Tanay, T., & Griffin, L. D. (2020). AI-enabled future crime. *Crime Science*, 9(1), 14. <https://doi.org/10.1186/s40163-020-00123-8>
- Carpenter, C. (2025). Whose [Crime] is it Anyway? *Journal of International Criminal Justice*, mqae055. <https://doi.org/10.1093/jicj/mqae055>



- Chen, N. (2025). Stolen Stories or Fair Use? *The New York Times v. OpenAI and the Limits of Machine Learning*. *Columbia Undergraduate Law Review*. <https://www.culawreview.org/ddc-x-culr-1/nyt-v-openai-and-microsoft>
- Christen, M., Burri, T., Kandul, S., & Vörös, P. (2023). Who is controlling whom? Reframing “meaningful human control” of AI systems in security. *Ethics and Information Technology*, 25(1), 10. <https://doi.org/10.1007/s10676-023-09686-x>
- Custers, B., Lahmann, H., & Scott, B. I. (2025). From liability gaps to liability overlaps: Shared responsibilities and fiduciary duties in AI and other complex technologies. *AI & SOCIETY*, 40(5), 4035–4050. <https://doi.org/10.1007/s00146-024-02137-1>
- Daniel, L. (2025, April 8). AI Avatars Replacing Human Lawyers In Court? Recent Case Says Not So Fast. *Forbes*. <https://www.forbes.com/sites/larsdaniel/2025/04/08/ai-avatars-replacing-human-lawyers-in-court-recent-case-says-not-so-fast/>
- Daniele, L. (2024). Incidental harm of the civilian in international humanitarian law and its *Contra Legem* antonyms in recent discourses on the laws of war. *Journal of Conflict and Security Law*, 29(1), 21–54. <https://doi.org/10.1093/jcsl/krae004>
- Dathathri, S., See, A., Ghaisas, S., Huang, P.-S., McAdam, R., Welbl, J., Bachani, V., Kaskasoli, A., Stanforth, R., Matejovicova, T., Hayes, J., Vyas, N., Merey, M. A., Brown-Cohen, J., Bunel, R., Balle, B., Cemgil, T., Ahmed, Z., Stacpoole, K., Kohli, P. (2024). Scalable watermarking for identifying large language model outputs. *Nature*, 634(8035), 818–823. <https://doi.org/10.1038/s41586-024-08025-4>
- Farooq, A., & De Vreese, C. (2025). Deciphering authenticity in the age of AI: How AI-generated disinformation images and AI detection tools influence judgements of authenticity. *AI & SOCIETY*. <https://doi.org/10.1007/s00146-025-02416-5>
- Fernandez-Basso, C., Gutiérrez-Batista, K., Gómez-Romero, J., Ruiz, M. D., & Martín-Bautista, M. J. (2025). An AI knowledge-based system for police assistance in crime investigation. *Expert Systems*, 42(1), e13524. <https://doi.org/10.1111/exsy.13524>
- Fidler, M. (2025). Fragmentation of International Cybercrime Law. *Utah Law Review*, 2025(3), 737–804. <https://dc.law.utah.edu/cgi/viewcontent.cgi?article=1413&context=ulr>
- Fine, A., Le, S., & K. Miller, M. (2023). Content Analysis of Judges’ Sentiments Toward Artificial Intelligence Risk Assessment Tools. *Criminology, Criminal Justice, Law & Society*, 24(2), 31–46. <https://scholasticahq.com/criminology-criminal-justice-law-society/>
- Gaeta, P. (2024). Who Acts When Autonomous Weapons Strike? *Journal of International Criminal Justice*, 21(5), 1033–1055. <https://doi.org/10.1093/jicj/mqae001>
- Guembe, B., Azeta, A., Misra, S., Osamor, V. C., Fernandez-Sanz, L., & Pospelova, V. (2022). The Emerging Threat of Ai-driven Cyber Attacks: A Review. *Applied Artificial Intelligence*, 36(1), 2037254. <https://doi.org/10.1080/08839514.2022.2037254>
- Härmand, K. (2023). AI Systems’ Impact on the Recognition of Foreign Judgements: The Case of Estonia. *Juridica International*, 32, 107–118. <https://doi.org/10.12697/JI.2023.32.09>
- Hyslip, T., & Pittman, J. (2015). A Survey of Botnet Detection Techniques by Command and Control Infrastructure. *Journal of Digital Forensics, Security and Law*. <https://doi.org/10.15394/jdfsl.2015.1195>
- Königs, P. (2022). Artificial intelligence and responsibility gaps: What is the problem? *Ethics and Information Technology*, 24(3), 36. <https://doi.org/10.1007/s10676-022-09643-0>
- Krishna, V. V. (2024). AI and contemporary challenges: The good, bad and the scary. *Journal of Open Innovation: Technology, Market, and Complexity*, 10(1), 100178. <https://doi.org/10.1016/j.joit-mc.2023.100178>



- L. Kimbrough, J. (2025). Developing Lawyering Skills in the Age of Artificial Intelligence: A Framework for Legal Education. *Journal of Technology Law & Police*, 29(1), 31–71. <https://scholarship.law.ufl.edu/cgi/viewcontent.cgi?article=1242&context=jtlp>
- Lake, T. (2024, April 26). A school principal faced threats after being accused of offensive language on a recording. Now police say it was a deepfake. *CNN*. <https://edition.cnn.com/2024/04/26/us/pikesville-principal-maryland-deepfake-cec/index.html>
- López-Borrull, A., & Lopezosa, C. (2025). Mapping the Impact of Generative AI on Disinformation: Insights from a Scoping Review. *Publications*, 13(3), 33. <https://doi.org/10.3390/publications13030033>
- Lundberg, E., & Mozelius, P. (2025). The potential effects of deepfakes on news media and entertainment. *AI & SOCIETY*, 40(4), 2159–2170. <https://doi.org/10.1007/s00146-024-02072-1>
- Mahardhika, V., Astuti, P., & Mustaffa, A. (2023). Could Artificial Intelligence be the Subject of Criminal Law? *Yustisia Jurnal Hukum*, 12(1), 1. <https://doi.org/10.20961/yustisia.v12i1.56065>
- Malik, A. W., Bhatti, D. S., Park, T.-J., Ishtiaq, H. U., Ryou, J.-C., & Kim, K.-I. (2024). Cloud Digital Forensics: Beyond Tools, Techniques, and Challenges. *Sensors*, 24(2), 433. <https://doi.org/10.3390/s24020433>
- Meurs, T., Cartwright, E., Cartwright, A., Junger, M., & Abhishta, A. (2024). Deception in double extortion ransomware attacks: An analysis of profitability and credibility. *Computers & Security*, 138, 103670. <https://doi.org/10.1016/j.cose.2023.103670>
- Miranda, B. (2025, August 5). ‘We didn’t vote for ChatGPT’: Swedish PM under fire for using AI in role. *The Guardian*. <https://www.theguardian.com/technology/2025/aug/05/chat-gpt-swedish-pm-ulf-kristersson-under-fire-for-using-ai-in-role>
- Mohsendokht, M., Li, H., Kontovas, C., Chang, C.-H., Qu, Z., & Yang, Z. (2024). Decoding dependencies among the risk factors influencing maritime cybersecurity: Lessons learned from historical incidents in the past two decades. *Ocean Engineering*, 312, 119078. <https://doi.org/10.1016/j.oceaneng.2024.119078>
- Montgomery, B. (2025, July 1). AI companies start winning the copyright fight. *The Guardian*. <https://www.theguardian.com/technology/2025/jun/30/ai-techscape-copyright>
- Nastoska, A., Jancheska, B., Rizinski, M., & Trajanov, D. (2025). Evaluating Trustworthiness in AI: Risks, Metrics, and Applications Across Industries. *Electronics*, 14(13), 2717. <https://doi.org/10.3390/electronics14132717>
- Nerantzi, E., & Sartor, G. (2024). ‘Hard AI Crime’: The Deterrence Turn. *Oxford Journal of Legal Studies*, 44(3), 673–701. <https://doi.org/10.1093/ojls/gqae018>
- Pantanowitz, L., Hanna, M., Pantanowitz, J., Lennerz, J., Henricks, W. H., Shen, P., Quinn, B., Bennet, S., & Rashidi, H. H. (2024). Regulatory Aspects of Artificial Intelligence and Machine Learning. *Modern Pathology*, 37(12), 100609. <https://doi.org/10.1016/j.modpat.2024.100609>
- Park, P. S., Goldstein, S., O’Gara, A., Chen, M., & Hendrycks, D. (2024). AI deception: A survey of examples, risks, and potential solutions. *Patterns*, 5(5), 100988. <https://doi.org/10.1016/j.pat-ter.2024.100988>
- Pesetski, A. (2020). Deepfakes: A New Content Category for a Digital Age. *WILLIAM & MARY BILL OF RIGHTS JOURNAL*, 29(2), 503–532. <https://scholarship.law.wm.edu/cgi/viewcontent.cgi?article=1965&context=wmborj>
- Popa, C., Pallath, R., Cunningham, L., Tahiri, H., Kesavarajah, A., & Wu, T. (2025). *Deepfake Technology Unveiled: The Commoditization of AI and Its Impact on Digital Trust* (arXiv:2506.07363). arXiv. <https://doi.org/10.48550/arXiv.2506.07363>
- Rahman-Jones, I. (2025, August 28). AI firm says its technology weaponised by hackers. *BBC*. <https://www.bbc.com/news/articles/crr24eqnnq9o>



- Runyon, N. (2025, May 8). Deepfakes on trial: How judges are navigating AI evidence authentication. *Thomson Reuters*. <https://www.thomsonreuters.com/en-us/posts/ai-in-courts/deepfakes-evidence-authentication/>
- Sandoval, M.-P., De Almeida Vau, M., Solaas, J., & Rodrigues, L. (2024). Threat of deepfakes to the criminal justice system: A systematic review. *Crime Science*, 13(1), 41. <https://doi.org/10.1186/s40163-024-00239-1>
- Sherman, N., & Hooker, L. (2025, June 25). Judge backs AI firm over use of copyrighted books. *BBC*. <https://www.bbc.com/news/articles/c77vr00enzyo>
- Steven, G. (2021). It's Time to Put Character Back into the Character-Evidence Rule. *Marquette Law Review*, 104(3), 709–811. <https://scholarship.law.marquette.edu/cgi/viewcontent.cgi?article=5483&context=mulr>
- Thomas, D. (2025, January 14). Judge rebukes Minnesota over AI errors in “deepfakes” lawsuit. *Reuters*. <https://www.reuters.com/legal/government/judge-rebukes-minnesota-over-ai-errors-deep-fakes-lawsuit-2025-01-13/>
- Wang, X. (2024). Global (re-)framing of cybercrime: An emerging common interest in flux of competing normative powers? *Leiden Journal of International Law*, 1–27. <https://doi.org/10.1017/S0922156524000402>
- Young, F. (2025). A Deepfake Evidentiary Rule (Just in Case). *University of Illinois Chicago*. <https://library.law.uic.edu/news-stories/a-deepfake-evidentiary-rule-just-in-case/>
- Završnik, A. (2020). Criminal justice, artificial intelligence systems, and human rights. *ERA Forum*, 20(4), 567–583. <https://doi.org/10.1007/s12027-020-00602-0>