

Embedding explainable AI into financial compliance education: Building capacity for cross-border transaction oversight

Meron Aberra Tassew* and Nneoma Azubuike

Marshall School of Business, University of Southern California, Los Angeles, United States.

Open Access Research Journal of Science and Technology, 2025, 15(01), 020-027

Publication history: Received on 10 August 2025; revised on 15 September 2025; accepted on 18 September 2025

Article DOI: <https://doi.org/10.53022/oarjst.2025.15.1.0114>

Abstract

The global financial system faces an escalating threat from sophisticated money laundering and illicit finance, facilitated by high-volume, cross-border transactions. Traditional compliance training, reliant on static case studies and rule memorisation, is critically inadequate for preparing professionals to oversee advanced, AI-driven surveillance systems, whose opaque “black box” nature erodes trust and accountability. The study aims to demonstrate how Explainable AI (XAI) can be systematically integrated into professional education to build essential human capacity for the effective oversight of algorithmic compliance tools. The proposed method involves the development of an XAI-infused learning architecture. This model incorporates a dynamic scenario engine, generating complex, real-world transaction simulations, and an XAI interrogation interface that allows trainees to actively query AI decisions using techniques like SHAP and counterfactual explanations. This is operationalised through the D.R.I.V.E.N. methodology, “Define, Review, Interrogate, Validate, Evaluate, Navigate”, a structured framework for critical engagement. The application of this model, illustrated through a Hawala network case study, shows a transformation in learning outcomes. Trainees evolve from passive recipients of alerts into active analysts capable of deconstructing AI reasoning. This fosters a critical, interrogative mindset, strengthening decision-making justification and embedding human oversight within the operational workflow. The integration of XAI into compliance education is an urgent strategic imperative. It is a necessary evolution to create a resilient human-machine partnership, mitigating algorithmic bias and enhancing the integrity of financial systems. Despite implementation challenges, this approach is vital for building a workforce capable of governing the future of automated finance.

Keywords: Explainable AI (XAI); Financial Compliance Education; Anti-Money Laundering; Regulatory Technology (RegTech); Pedagogical Model; Cross-Border Transactions; Algorithmic Oversight

1. Introduction

The global financial system is defined by its increasing velocity and complexity. Cross-border capital flows, although essential for economic growth, pose a significant challenge for those responsible for maintaining the system’s integrity (Chang, Iakovou, & Shi, 2020). The scale is staggering; trillions of dollars move across jurisdictions daily, creating a vast landscape within which illicit finance can be concealed (Financial Action Task Force, 2023). Money launderers and providers of corruption financing continuously refine their methods, employing sophisticated techniques that exploit gaps in national oversight and the limitations of conventional monitoring tools (Gaviyau & Sibindi, 2023). This escalating threat landscape renders traditional compliance training programmes increasingly inadequate. These programmes, often reliant on historical case studies and static rule memorisation, fail to equip professionals with the necessary skills to investigate modern, algorithm-driven financial crime (Alsuwailam & Saudaga, 2020).

Consequently, regulatory technology or RegTech has emerged as a necessary response. Financial institutions now deploy artificial intelligence and machine learning systems to automate transaction monitoring, screen clients against sanctions lists, and identify anomalous patterns indicative of malfeasance (Aziz & Andriansyah, 2023). These systems

* Corresponding author: Meron Aberra Tassew

offer a powerful advantage in processing speed and analytical depth. However, their adoption introduces a significant pedagogical and operational dilemma. The inherent opacity of many advanced algorithms, frequently described as black boxes, creates a critical deficit in understanding. A compliance officer may receive an alert from a system but lack any clear insight into the reasoning behind it (Martinez, 2020). This lack of explainability erodes trust, complicates decision-making, and poses serious problems for accountability, particularly when a justification for a suspicious activity report must be presented to regulators or courts.

This gap between technological capability and human understanding forms the central problem this research addresses. The question is not whether to use AI in compliance, but how to educate the human agents, compliance officers, regulators, and policymakers to work alongside it effectively. There is an urgent need for a new pedagogical model that moves beyond teaching rules and towards cultivating a deep, intuitive comprehension of AI-driven processes. This requires embedding the principles and practical applications of Explainable AI directly into professional education. This study proposes such a model. It argues for the integration of interactive XAI systems into training frameworks, enabling professionals to interrogate, understand, and validate AI outputs. The objective is to build a crucial human capacity for oversight, thereby strengthening the entire defence against financial crime in cross-border environments.

2. The convergence of three domains

The proposed model for integrating Explainable AI into compliance education exists at the intersection of three distinct academic and professional fields. A critical examination of each reveals a terrain marked by stagnation, over-promising, and a concerning lack of cross-disciplinary dialogue. The first domain, financial compliance training, remains largely anchored in an outdated pedagogical paradigm. For decades, training has been dominated by passive knowledge transfer, focusing on the rote memorisation of regulations and the retrospective analysis of historical case studies (Power, 2013). This approach produces technicians who can recite rules but who are ill-equipped for the proactive, analytical, and rapidly evolving nature of modern financial crime. The recent shift towards digital learning platforms, while a welcome modernisation, often amounts to little more than placing these same outdated materials online, mistaking technological delivery for pedagogical innovation (Timotheou et al., 2023; Bitar & Davidovich, 2024). This field has failed to keep pace with the very crimes it seeks to prevent, creating a workforce that is perpetually behind the curve.

At the same time, the field of artificial intelligence in finance has surged forward, but with a pronounced and problematic bias towards technical optimisation over practical applicability. A significant body of literature enthusiastically documents the development of increasingly sophisticated machine learning models for fraud detection, network analysis, and anomaly monitoring (Vuković, Dekpo-Adza, and Matović, 2025). The dominant metrics for success are computational efficiency and predictive accuracy. However, this technical pursuit often wilfully ignores the operational context in which these tools must function. The result is a proliferation of so-called “black box” systems whose decision-making processes are inaccessible to the human agents tasked with acting upon their outputs. This creates a dangerous accountability gap. As Ananny and Crawford (2018) and Montagnani, Najjar and Davola (2024) argue, opacity in automated systems shifts burdens of responsibility and undermines the possibility of meaningful oversight. In the high-stakes realm of financial compliance, where decisions can freeze assets and ruin reputations, an unexplained algorithm output is not a solution; it is a profound liability.

The emerging field of Explainable AI (XAI) presents itself as an antidote to this opacity, yet it too suffers from a critical limitation. Research in XAI, including work on techniques like LIME and SHAP (Ribeiro, Singh, and Guestrin, 2016), has been predominantly technocentric. Its focus is on making the machine understandable to other machines or to computer scientists, not on making it comprehensible and useful for domain experts like compliance officers. There is a mistaken assumption that providing a technical explanation, such as a feature importance score, equates to fostering genuine understanding in a non-technical professional. This overlooks the vast chasm between algorithmic explanation and human reasoning, a chasm filled with contextual knowledge, professional judgment, and ethical considerations. The work of Miller (2019) on the social and psychological dimensions of explanation highlights that effectiveness is not a technical property but is rather contingent on the recipient’s knowledge and goals. XAI research has largely failed to incorporate these insights from cognitive science and the learning sciences, limiting its real-world utility.

The convergence of these three domains reveals a critical scholarly and practical gap. The compliance training literature lacks technological sophistication, the AI-in-finance literature lacks pedagogical and operational awareness, and the XAI literature lacks a user-centred design focus grounded in educational theory. They operate in silos, publishing in their own journals and attending their own conferences, with minimal cross-pollination. The consequence is that the human element, the compliance officer who must interpret, question, and act, has become the weakest link in the chain. A new approach is urgently needed, one that does not merely layer these domains but truly integrates them, creating a new

pedagogical science designed for the age of automated finance. This requires moving beyond technical explanations to foster a critical, interrogative, and deeply contextual form of literacy that allows professionals to effectively govern the machines upon which they increasingly rely.

3. The proposed model: An XAI-infused learning architecture

Current approaches to compliance training are fundamentally inadequate for the demands of modern finance. They represent a pedagogical failure, treating the symptom, a lack of rule knowledge, while ignoring the disease, a critical deficit in analytical reasoning and technological literacy. The model proposed here is designed to address this core failure. It is not an incremental improvement but a foundational shift towards a learning architecture that privileges critical interrogation over passive receipt of information. This architecture is built upon a specific philosophical stance: that effective oversight in the age of AI is not achieved by training humans to trust algorithms, but by equipping them to critically distrust and systematically verify algorithmic outputs. This pedagogy of scepticism is the essential defence against automated bias and the abdication of professional judgment (Pasquale, 2015).

The model comprises several integrated components, each designed to counteract a specific deficiency in traditional training. The first is a Dynamic Scenario Engine. This system moves beyond static, pre-written case studies by generating a near-infinite variety of simulated cross-border transactions. These scenarios are not mere illustrations; they are complex, data-rich environments that mimic the overwhelming noise and ambiguous signals of a real-world compliance officer's dashboard. The engine parameters can be adjusted to reflect specific typologies, such as trade-based money laundering or the layering of illicit funds through cryptocurrency mixers, based on real-world data from institutions like the Financial Conduct Authority (2025).

The second, and most critical, component is the XAI Interrogation Interface. This is the practical manifestation of the pedagogy of scepticism. When the trainee, acting as a compliance officer, receives an alert on a transaction or client, they cannot simply proceed. The system mandates an active process of interrogation. The interface provides a suite of tools, each mapped to a specific XAI technique. A trainee can select a 'Feature Importance' query, which would call upon a SHAP-based analysis to visually rank the factors that most influenced the model's decision, such as transaction velocity or jurisdiction risk. Alternatively, a 'Counterfactual Explanation' tool would generate a statement: "This payment would not have been flagged if the beneficiary country had a FATF compliance score above 4." This forces the trainee to consider the precise logic thresholds of the model. A third tool, the 'Adversarial Example' generator, allows the trainee to attempt to manipulate transaction parameters to bypass the AI's detection, revealing the model's potential blind spots and strengthening the trainee's understanding of its limitations. This process transforms the AI from an oracle to be believed into a colleague to be questioned.

The third component is an Adaptive Learning Pathway. The system continuously assesses trainee performance, but not merely on whether they correctly flag a transaction. The primary metric of success is the quality of the trainee's justification for their decision, based on the evidence generated through the XAI Interrogation Interface. A trainee who consistently accepts AI alerts without independent analysis will be progressively served scenarios where the AI model has known, embedded biases or high false-positive rates. On the other hand, a trainee who excessively dismisses accurate alerts will face scenarios demonstrating the severe consequences of missing sophisticated illicit flows. This adaptive system ensures that the learning experience is personalised to challenge cognitive biases and build robust decision-making heuristics. A comparative overview of the pedagogical shift is presented in Table 1.

Table 1 Contrasting Pedagogical Approaches to Compliance Training

Aspect	Traditional Training Model	XAI-Infused Learning Architecture
Core Objective	Knowledge transfer of rules and historical cases.	Development of critical analytical skill to interrogate AI systems.
Case Studies	Static, predetermined, and simplified.	Dynamic, generative, complex, and data-rich.
Role of AI	Absent or presented as an unproblematic solution.	Central and presented as a tool requiring verification.
Trainee Role	Passive recipient of information.	Active investigator and sceptical interrogator.
Success Metric	Correct classification of a scenario.	Quality of justification using explanatory evidence.

Adaptivity	One-size-fits-all content delivery.	Pathway dynamically adjusts to counter individual biases.
------------	-------------------------------------	---

Finally, the entire architecture is framed within a rigorous regulatory context. Every scenario and AI model is explicitly mapped against the relevant provisions of global frameworks, such as the FATF Recommendations or the EU's Sixth Anti-Money Laundering Directive. This ensures that the exercise of technological interrogation is always grounded in a specific legal and regulatory purpose. The trainee is not just learning how an algorithm works; they are learning how to govern its application within a specific rule of law, ensuring that technological efficiency never supersedes legal necessity.

This model does not propose a simple technical fix. It demands a significant investment in content development, computational resources, and instructional design. Its implementation would reveal the profound weaknesses in our current preparation of financial professionals. However, the cost of not making this investment is a compliance regime that is increasingly automated yet increasingly blind, managed by professionals who are present yet increasingly irrelevant. This architecture is proposed as a necessary starting point for building a workforce capable of governing the machines that will inevitably shape the future of finance.

4. A practical framework: The D.R.I.V.E.N. methodology for XAI training

The theoretical architecture demands a practical and operationalisable methodology. Without a clear procedural framework, the risk exists that this new pedagogy would degenerate into a disconnected series of technical exercises, failing to instil the disciplined thought processes required for high-stakes oversight. The D.R.I.V.E.N. methodology is proposed to structure the trainee's interaction within the XAI learning environment. This acronym encapsulates a cyclic process of inquiry and validation, designed to systematically dismantle the opacity of AI decision-making and rebuild it through human-led verification. It is a scaffold for cultivating what might be termed 'algorithmic scepticism', a necessary professional disposition for the 21st century (O'Neil, 2017).

The methodology begins with "Define". The trainee must first establish the investigative parameters. This involves selecting the appropriate AI model profile, perhaps one optimised for detecting trade-based money laundering versus one for spotting sanctioned entity evasion, and configuring the initial risk rules based on the simulated institution's risk appetite. This step counteracts the common tendency to treat AI as a monolithic oracle, forcing an understanding that different algorithms possess different strengths, weaknesses, and intended uses.

The subsequent phase, "Review", involves the initial receipt of the AI's output, typically an alert or a risk score. Critically, this output is not taken as a conclusion but as a hypothesis requiring immediate and rigorous testing. The core of the methodology is the "Interrogate Stage". Here, the trainee actively deploys the XAI tools described previously. They must generate and interpret SHAP value plots to see which features drove the model's decision, demand counterfactual explanations to understand the model's decision boundaries, and analyse local surrogate models to approximate the complex AI's behaviour in a more interpretable way. This is not a passive review; it is an active cross-examination.

Following interrogation, the trainee must "Validate". This is the crucial point of human judgment. The trainee synthesises the explanations provided by the XAI system with their own domain knowledge, contextual understanding of the client, and awareness of current illicit finance typologies. They must make a binary decision: to validate the AI's alert and proceed to escalation, or to reject it with a justified override. The system then provides an "Evaluate feedback loop", scoring the decision not on its correctness per se, but on the evidential quality and logical coherence of the justification provided. Finally, the trainee learns to "Navigate" the correct procedural aftermath of their decision, whether that involves drafting a suspicious activity report or documenting the reasons for dismissal, thereby closing the loop on a complete compliance action. Table 2 below shows a glimpse of the methodology in practice.

Table 2 The D.R.I.V.E.N. Methodology in Practice

Stage	Trainee Action	Pedagogical Objective
Define	Configures the AI model and rule-set parameters.	Understands that AI systems are tools with configurable purposes, not neutral arbiters.
Review	Receives and parses the initial AI-generated alert.	Learns to treat algorithmic outputs as hypotheses, not conclusions.

Interrogate	Actively queries the AI using XAI tools (e.g., SHAP, counterfactuals).	Develops skills to deconstruct and critically assess the logic behind AI decisions.
Validate	Makes a final decision based on XAI evidence and personal judgment.	Cultivates the authority to agree or override the AI, reinforcing human agency.
Evaluate	Receives feedback on the rationale behind their decision.	Strengthens metacognitive skills and the ability to articulate evidential reasoning.
Navigate	Executes the correct procedural steps following their decision.	Embeds the analytical exercise within a full operational workflow.

This structured approach moves training beyond button-clicking towards genuine cognitive apprenticeship. It acknowledges that true expertise lies not in blindly following a system's recommendations, but in possessing the critical faculties to understand and, when necessary, challenge them. However, the ultimate test of this framework lies in its application to a realistic and high-stakes scenario, demonstrating its capacity to build resilience against evolving threats. The following case study applies the D.R.I.V.E.N. methodology to a notoriously challenging area of compliance.

5. Case study illustration: Applying the model to a hawala transaction network

The true measure of any pedagogical framework is its performance when confronted with a notoriously opaque and challenging real-world phenomenon (Windschitl, 2002). The informal value transfer system known as Hawala presents a perfect stress test for the proposed XAI training model. Traditional training methods often reduce Hawala to a simplistic definition, failing to convey the complex transactional patterns that obscure its use for illicit purposes (Siddiqui, 2014; Soudjin, 2015). A trainee educated through conventional means would likely miss the subtle signals amidst legitimate financial flows. This case study illustrates the application of the D.R.I.V.E.N. methodology to a simulated Hawala network, demonstrating a path towards more effective detection.

A trainee is presented with a client, a small electronics importer, who has just received a series of payments from a counterparty in a high-risk jurisdiction. The AI system generates a high-risk alert. A traditionally trained officer might see ostensibly legitimate trade transactions and dismiss the alert, representing a critical failure. Applying the D.R.I.V.E.N. methodology, the trainee first "Defines" the scenario, selecting an AI model profile tuned for detecting trade-based money laundering and setting a low threshold for anomaly detection given the jurisdictional risk.

Upon "Reviewing" the alert, the trainee notes the AI's initial risk score is 94%. Instead of accepting this figure, the trainee moves to "Interrogate". Using the XAI interface, the trainee requests a feature importance explanation. A SHAP value plot reveals that the transaction amounts are the primary driver, consistently hovering just below the client's usual reporting threshold. A counterfactual explanation further indicates that the alert would not have triggered if the payments had been consolidated into larger, less frequent transfers, suggesting deliberate structuring.

Table 3 XAI Interrogation of a Simulated Hawala Alert

XAI Tool Used	Output Provided	Trainee Interpretation
Feature Importance (SHAP)	Identifies transaction amount and frequency as top contributors.	Recognises potential structuring behaviour to avoid detection thresholds.
Counterfactual Explanation	States alert would not trigger with fewer, larger payments.	Confirms hypothesis of deliberate transaction patterning.
Network Analysis	Reveals indirect links between sender and entities with prior SARs.	Identifies a potential node within a broader network of concern, beyond the immediate transaction.

The trainee then uses a network analysis tool, which visualises the transaction pathways. This reveals that the sending counterparty, while a registered business, is connected to several other entities that have been the subject of recent Suspicious Activity Reports for unrelated activities. This contextual network information, invisible in the raw transaction data, is critical. The trainee must now "Validate". They synthesise the XAI evidence, the structured amounts, the high-risk jurisdiction, and the network connections, with their knowledge of Hawala's reliance on net settlement across jurisdictions. They validate the alert and proceed to "Navigate" the process of filing a suspicious activity report,

meticulously documenting each piece of explanatory evidence provided by the XAI system. Table 3 reveals the XAI Interrogation of a Simulated Hawala Alert.

This exercise moves the trainee beyond a binary decision on a single transaction. It forces an investigation into the behavioural and relational patterns characteristic of complex money laundering schemes. The trainee learns that their role is not to blindly trust the AI's score but to forensically reconstruct the reasoning behind it, using that information to build an extensive case. This depth of analysis, which seamlessly integrates algorithmic output with human expertise, represents a significant advance over current training outcomes. The implications of embedding such critical thinking skills across the financial compliance workforce are substantial and merit further discussion.

6. Implications for stakeholders and financial stability

The implementation of an XAI-infused educational model would fundamentally alter the operational dynamics of financial oversight, with profound and critical implications for all stakeholders. For compliance officers, this represents a necessary but disruptive evolution of their professional identity. The role would transition from that of a passive rule-checker to an active forensic analyst and AI supervisor. This shift demands a higher order of critical thinking and technological literacy, raising serious questions about the current skills gap within the profession (Morandini et al., 2023). Without substantial investment in continuous professional development of this nature, compliance officers risk becoming irrelevant, merely rubber-stamping the decisions of opaque algorithms they cannot comprehend, thus creating a new and dangerous form of automated bias (Zuboff, 2019).

For regulators, this model offers a potential pathway to more effective supervision but also exposes a significant capability deficit. A regulator trained in XAI interrogation would be equipped to audit an institution's AI systems not just for outcomes, but for the robustness and fairness of their internal decision-making processes (Rane, Choudhary, and Rane, 2023). They could move beyond asking "what does your model do?" to demanding "show me how and why it reached this specific conclusion." This would elevate regulatory examinations from a compliance checklist to a substantive review of algorithmic governance. However, it presupposes that regulatory bodies can attract and retain talent with the requisite expertise to keep pace with the private sector, a challenge most are currently failing to meet.

At a systemic level, the widespread adoption of such training could enhance financial stability. A workforce capable of critically overseeing AI systems would form a more resilient and adaptive first line of defence against financial crime (Jada and Mayayise, 2024). This human-centric layer of oversight acts as a crucial circuit-breaker, mitigating the inherent model risk and potential procyclicality of automated systems that might otherwise amplify errors across the network (Danielsson, 2013). It moves the system away from a fragile reliance on imperfect algorithms towards a more robust human-machine partnership where each component checks the other's weaknesses. However, these theoretical benefits are contingent upon overcoming significant practical obstacles that threaten to undermine the entire endeavour.

7. Challenges and limitations

Notwithstanding its potential, the proposed model faces formidable practical and conceptual obstacles that must be acknowledged. A primary concern is the foundational issue of data integrity. The pedagogical effectiveness of any XAI system is wholly dependent on the quality and representativeness of its training data. An AI model trained on historically biased data will produce biased outputs, and consequently, its explanations will rationalise and perpetuate those same biases (Nishant, Schneckenberg, & Ravishankar, 2023; Ferrara, 2024). Training a professional to expertly interrogate a flawed system risks lending a veneer of legitimacy to discriminatory outcomes, a deeply troubling prospect for equitable enforcement.

Furthermore, the significant resource requirements present a major barrier to adoption. Developing and maintaining a dynamic, AI-driven learning platform with realistic financial data simulations demands substantial investment in computational infrastructure, data licensing, and specialised instructional design expertise. This cost could prove prohibitive for smaller regulatory bodies and educational institutions, potentially creating a two-tier system where only the largest, wealthiest organisations can develop this crucial capacity. This compounds existing inequalities in global financial oversight.

In addition, resistance to change within established institutions also poses a significant risk. The model's success hinges on a cultural shift that values critical interrogation over procedural compliance. Seasoned professionals and traditional training departments may view this new approach with scepticism, perceiving it as an unnecessary complication or a

challenge to established expertise. Overcoming this institutional inertia requires demonstrating not only the model's efficacy but also its necessity in an evolving financial landscape. Finally, the model itself must be continuously updated to reflect emerging criminal methodologies and evolving XAI techniques, creating a perpetual cycle of development that demands ongoing commitment. These substantial hurdles must be comprehensively addressed before the full benefits of such an educational transformation can be realised.

8. Conclusion

This research has argued that the escalating sophistication of financial crime necessitates a fundamental reconceptualisation of compliance education. Traditional training methods, focused on static rule memorisation, are largely inadequate for an era defined by algorithmic oversight. The proposed model, integrating Explainable AI into a dynamic, interrogative learning architecture, represents a necessary pedagogical shift. It moves the focus from passive knowledge receipt to active critical engagement, fostering a workforce capable of governing complex AI systems rather than being subservient to them. The D.R.I.V.E.N. methodology provides a practical framework for cultivating this essential algorithmic scepticism. While significant implementation challenges exist, from data bias to institutional resistance, the alternative, a compliance regime managed by professionals who cannot understand the tools they use, poses a far greater risk to the integrity of the global financial system. Building human capacity for oversight is not a technical luxury but an urgent strategic imperative for safeguarding financial stability.

Compliance with ethical standards

Disclosure of conflict of interest

No conflict of interest to be disclosed.

References

- [1] Alsuwailam, A.A.S. and Saudagar, A.K.J. (2020). Anti-money laundering systems: a systematic literature review. *Journal of Money Laundering Control*, 23(4), pp.833-848.
- [2] Ananny, M. and Crawford, K. (2018). Seeing Without Knowing: Limitations of the Transparency Ideal and Its Application to Algorithmic Accountability. *New Media & Society*, 20, pp.973-989. <https://doi.org/10.1177/1461444816676645>
- [3] Aziz, L.A.R. and Andriansyah, Y. (2023). The role artificial intelligence in modern banking: an exploration of AI-driven approaches for enhanced fraud prevention, risk management, and regulatory compliance. *Reviews of Contemporary Business Analytics*, 6(1), pp.110-132.
- [4] Bitar, N. and Davidovich, N. (2024). Transforming Pedagogy: The Digital Revolution in Higher Education. *Education Sciences*, 14(8), 811. <https://doi.org/10.3390/educsci14080811>
- [5] Chang, Y., Iakovou, E. and Shi, W. (2020). Blockchain in global supply chains and cross border trade: a critical synthesis of the state-of-the-art, challenges and opportunities. *International Journal of Production Research*, 58(7), pp.2082-2099.
- [6] Danielson, C. (2013). *The Framework for Teaching: Evaluation Instrument*. Princeton, NJ: The Danielson Group.
- [7] Ferrara, E. (2024). Fairness and Bias in Artificial Intelligence: A Brief Survey of Sources, Impacts, and Mitigation Strategies. *Sci*, 6(1), p.3. <https://doi.org/10.3390/sci6010003>
- [8] Financial Action Task Force (2023). *Illicit Financial Flows from Cyber-Enabled Fraud*. Available at: <https://www.fatf-gafi.org/content/dam/fatf-gafi/reports/Illicit-financial-flows-cyber-enabled-fraud.pdf.coredownload.inline.pdf>
- [9] Financial Conduct Authority (2025). *Financial Crime Thematic Reviews*. Available at: <https://www.handbook.fca.org.uk/handbook/FCTR.pdf>
- [10] Gaviyau, W. and Sibindi, A.B. (2023). Global Anti-Money Laundering and Combating Terrorism Financing Regulatory Framework: A Critique. *Journal of Risk and Financial Management*, 16(7), 313. <https://doi.org/10.3390/jrfm16070313>

- [11] Jada, I. and Mayayise, T.O. (2024). The impact of artificial intelligence on organisational cyber security: An outcome of a systematic literature review. *Data and Information Management*, 8(2), p.100063. <https://doi.org/10.1016/j.dim.2023.100063>
- [12] Martinez, V.R. (2020). Complex compliance investigations. *Columbia Law Review*, 120(2), pp.249-308.
- [13] Miller, T. (2019). Explanation in artificial intelligence: insights from the social sciences. *Artificial Intelligence*, 267, pp.1-38. <https://doi.org/10.1016/j.artint.2018.07.007>
- [14] Montagnani, M.L., Najjar, M.C. and Davola, A. (2024). The EU Regulatory approach(es) to AI liability, and its Application to the financial services market. *Computer Law & Security Review*, 53, p.105984.
- [15] Morandini, S., Fraboni, F., De Angelis, M., Puzzo, G., Giusino, D. and Pietrantonio, L. (2023). The impact of artificial intelligence on workers' skills: Upskilling and reskilling in organisations. *Informing Science*, 26, pp.39-68. <https://doi.org/10.28945/5078>
- [16] Nishant, R., Schneckenberg, D. and Ravishankar, M. (2023). The formal rationality of artificial intelligence-based algorithms and the problem of bias. *Journal of Information Technology*, 39(1), pp.19-40. <https://doi.org/10.1177/02683962231176842>
- [17] O'neil, C. (2017). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown.
- [18] Pasquale, F. (2015). *The black box society: The secret algorithms that control money and information*. Harvard University Press.
- [19] Power, M. (2013). The apparatus of fraud risk. *Accounting, Organizations and Society*, 38(6-7), pp.525-543.
- [20] Rane, N., Choudhary, S. and Rane, J. (2023). Explainable Artificial Intelligence (XAI) approaches for transparency and accountability in financial decision-making. SSRN. Available at: <https://ssrn.com/abstract=4640316>
- [21] Ribeiro, M.T., Singh, S. and Guestrin, C. (2016). "Why Should I Trust You?" Explaining the Predictions of Any Classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, San Francisco, 13-17 August 2016, pp.1135-1144. <https://doi.org/10.1145/2939672.2939778>
- [22] Siddiqui, S.K. (2014). *The regulation of Hawala and other IVTS in Post 9/11 Years: A case study of Pakistan's Hawala Regulation 2002*. Doctorate Dissertation, University of Wollongong.
- [23] Soudjin, M. (2015). Hawala and money laundering: Potential use of red flags for persons offering hawala services. *European Journal on Criminal Policy and Research*, 21(2), pp.257-274.
- [24] Timotheou, S., Miliou, O., Dimitriadis, Y., Sobrino, S.V., Giannoutsou, N., Cachia, R., Monés, A.M. and Ioannou, A. (2023). Impacts of digital technologies on education and factors influencing schools' digital capacity and transformation: A literature review. *Education and Information Technologies*, 28(6), pp.6695-6726. <https://doi.org/10.1007/s10639-022-11431-8>
- [25] Vuković, D.B., Dekpo-Adza, S. and Matović, S. (2025). AI integration in financial services: a systematic review of trends and regulatory challenges. *Humanities and Social Sciences Communications*, 12(1), pp.1-29. <https://doi.org/10.1057/s41599-025-04850-8>
- [26] Windschitl, M. (2002). Framing constructivism in practice as the negotiation of dilemmas: An analysis of the conceptual, pedagogical, cultural, and political challenges facing teachers. *Review of Educational Research*, 72(2), pp.131-175.
- [27] Zuboff, S. (2019). *Surveillance Capitalism and the Challenge of Collective Action*. *New Labor Forum*, 28(1), pp.10-29. <https://doi.org/10.1177/1095796018819461>.