

EL-RAKHAWI THEORY OF PROACTIVE JURISPRUDENCE FOR ARTIFICIAL INTELLIGENCE

From Reactive Accountability to Predictive Governance in the Age of Complex Intelligent Systems

A Complete Legal, Ethical, Predictive, and Technological Framework for Building Proactive Accountability Systems Before the Occurrence of Crises

Author: Dr. Mohamed Kamal Arafa Elrakhawi

Legal Author, International Lecturer in Law, Jurist, Expert, Researcher, and Consultant

DOI: 10.5281/zenodo.21868244

Date of Publication: July 2026

Languages: Arabic, English

Status: A Theoretical, Legal, Ethical, Predictive, and Technological Academic Framework

DEDICATION

To the spirit of my father and the soul of my mother, may God rest them in peace. To the two trusts of the two worlds: knowledge is entrusted with justice, and science is entrusted with wisdom. Justice is the foundation of sovereignty. Repelling corruption takes precedence over bringing benefit. And prevention is better than cure.

And to my two beloved daughters, the two dear eyes of Algeria and Egypt, Sabrina and Lina. And to every person who believes that technology must serve humanity, not enslave the instinct. And to every new generation that chooses freedom of consciousness. And to every human who refuses to be enslaved by the algorithm.

This work is the fruit of years of engagement with the problematic of accountability in the age of artificial intelligence. I hope it will be a building block in constructing a safer and more just future for science, law, and humanity.

EXECUTIVE SUMMARY

This reference presents an entirely new theoretical framework addressing a fundamental problematic in the age of artificial intelligence: How do we build strong and responsible accountability systems for artificial intelligence before the occurrence of crises, instead of waiting until disasters happen?

Current theories of legal accountability in artificial intelligence are based on the principle of reactive accountability: we wait until harm occurs, then we search for the responsible party, then

we apply the law. But in the age of generative artificial intelligence, autonomous agents, and complex systems, this approach may be catastrophic, irreparable, and crossing borders in seconds.

The Elrakhawi Theory of Proactive Jurisprudence for Artificial Intelligence presents a methodological solution to this problematic through eighteen innovative principles, presenting a radical shift from reactive accountability to predictive governance, where legal frameworks are built before the emergence of technologies, not after them:

Principle One: Proactive Jurisprudence. The principle that presents a radical shift from reactive accountability to predictive governance, where legal frameworks are built before the emergence of technologies, not after them.

Principle Two: The Model of Predicting Algorithmic Risks. The principle that presents a model for predicting algorithmic risks before their occurrence, based on the analysis of patterns and anomalies.

Principle Three: Proactive Safety Certificates. The principle that obligates developers of artificial intelligence to obtain safety certificates before deployment, proving their freedom from catastrophic risks.

Principle Four: The Distributed Algorithmic Accountability Registry. The principle that presents a blockchain-based registry that records every algorithmic decision in a decentralized manner, providing absolute transparency and permanent accountability.

Principle Five: The Algorithmic Emergency Protocol. The principle that presents automatic circuit breaker mechanisms that automatically stop dangerous systems upon discovery of risks.

Principle Six: The Proactive Ethics Committee. The principle that presents an independent committee to evaluate the ethical impact of artificial intelligence before its deployment.

Principle Seven: The Principle of Digital Precaution. The principle derived from Islamic jurisprudence that decides that in cases of doubt about system safety, it must not be deployed until its safety is proven.

Principle Eight: The Right to Digital Safety. The principle that presents legal recognition of the right of every human to protection from algorithmic harms.

Principle Nine: The Accurate Prediction Parameters. The principle that addresses the problematic of false predictions through the mechanism of prediction appeal and the addition of the prediction insurance fund.

Principle Ten: The Innovation Support Fund. The principle that addresses the problematic of stifling innovation through providing grants for small and emerging companies.

Principle Eleven: The Ethical Dilemmas Protocol. The principle that addresses the problematic of ethical dilemmas through multiple ethics committees with diverse specializations.

Principle Twelve: Anti-Repeated Bias Algorithms. The principle that addresses the problematic of bias in prediction algorithms through external periodic review.

Principle Thirteen: The Multi-Level Governance Model. The principle that addresses the problematic of technological sovereignty through governance at national, regional, and international levels.

Principle Fourteen: The Dynamic Adaptation Mechanism. The principle that addresses the problematic of rapid development through periodic review every six months.

Principle Fifteen: The Principle of Predictive Proportionality. The principle that addresses the problematic of responsibility for predictions through not punishing false predictions except in cases of gross negligence.

Principle Sixteen: The Super Governance Protocol. The principle that addresses the problematic of general artificial intelligence through the principle of the human in the loop.

Principle Seventeen: The Mixed Financing Model. The principle that addresses the problematic of economic cost through distributing financing among countries and companies.

Principle Eighteen: The Graduated Transparency Model. The principle that addresses the problematic of privacy and transparency through multiple levels of transparency.

The reference also presents a complete regulatory framework including an advanced mathematical model for risk prediction that includes dynamic prediction parameters, network parameters, and epistemic uncertainty parameters, along with anomaly detection algorithms, the cognitive robot, the distributed accountability registry, the proactive ethics committee, emergency protocols, and proactive safety certificates systems.

The reference relies on documented real cases including the Flash Crash of 2010, self-driving car accidents, bias in recruitment algorithms, and the risks of generative artificial intelligence. It presents a deep legal, ethical, and predictive analysis of how to apply the theory to these cases.

INTRODUCTION

On the sixth of May 2010, the American financial markets witnessed what became known as the Flash Crash, where the Dow Jones index lost approximately one trillion dollars in minutes. Dozens of algorithms of automated trading interacted with each other in a complex interaction. There was no single perpetrator behind this disaster.

In March 2018, a self-driving car developed by Uber killed a woman in Arizona. There was no human driver. The algorithm could have predicted the pedestrian in the appropriate time but did not.

In 2024, generative artificial intelligence models produced misleading images and texts that affected elections in several countries, and spread false medical information that caused deaths.

In all these cases, the legal response was reactive: we wait until harm occurs, then we search for the responsible party, then we apply the law. But this approach may be catastrophic, irreparable, and crossing borders in seconds, where harms occur at a speed and acceleration that exceed human capacity to respond.

This problematic is not new. Since the emergence of the internet, there have been growing calls to regulate the digital space. But regulation has always been reactive, with no anticipation.

In 2024, the European Artificial Intelligence Act entered into force, which is the first comprehensive law for artificial intelligence in the world. But it still relies on an approach of classifying risks after their occurrence, not predicting them before they happen.

This reference presents an entirely new theory that addresses this problematic in a deep and methodological manner: The Elrakhawi Theory of Proactive Jurisprudence for Artificial Intelligence. The theory is based on the radical shift from reactive accountability to predictive governance, and presents a complete legal, ethical, predictive, and technological framework for building proactive accountability systems before the occurrence of crises.

PART ONE

THE THEORETICAL, HISTORICAL, AND PHILOSOPHICAL FRAMEWORK

VOLUME ONE

THE EVOLUTION OF THE CONCEPT OF ACCOUNTABILITY IN THE DIGITAL AGE

CHAPTER ONE

FROM REACTIVE TO PROACTIVE ACCOUNTABILITY

In the industrial age, accountability was reactive: we wait until harm occurs, then we pursue the responsible party. Because harms were limited and reparable, this approach was sufficient.

But in the digital age, new problematics emerged. The internet created a space where harms could spread rapidly. Accountability remained reactive in the digital age, but it was no longer sufficient.

The point of transformation came in the age of artificial intelligence, where harms could be catastrophic, irreparable, and crossing borders in seconds. Reactive accountability is no longer sufficient.

CHAPTER TWO THE FAILURE OF TRADITIONAL MODELS

Traditional models of accountability in artificial intelligence rely on contractual liability, where we wait until one party enters the contract to apply liability. And strict liability, where we apply liability even without fault. And fault-based liability, where we wait until harm occurs then pursue the responsible party.

But these models face fundamental problematics in the age of artificial intelligence: difficulty of identifying the responsible party in algorithmic networks, the speed of harm occurrence that exceeds the capacity of law to respond, difficulty of proof, and the complexity of causal response in complex systems.

CHAPTER THREE THE NEED FOR PROACTIVE JURISPRUDENCE

Proactive jurisprudence is a new approach that builds legal frameworks before the emergence of technologies, not after them. It is based on predicting potential risks through the use of artificial intelligence, proactive prevention through building mechanisms of prevention before crises, and dynamic adaptation through updating frameworks in line with technological development.

This approach is not entirely new. In Islamic jurisprudence, the principles of blocking the means and repelling corruption represent forms of proactive jurisprudence.

VOLUME TWO ISLAMIC JURISPRUDENCE FOR PROACTIVE JURISPRUDENCE

CHAPTER FOUR RULES OF PRECAUTION AND PREVENTION IN ISLAMIC JURISPRUDENCE

Islamic jurisprudence provides a solid foundation for proactive jurisprudence through the rules of precaution and prevention.

The rule of repelling corruption takes precedence over bringing benefit decides that if corruption and benefit conflict, repelling corruption must come first. In the context of artificial intelligence: if the system carries potential catastrophic risks, it must not be deployed until its safety is proven conclusively.

The rule of blocking the means decides that paths leading to prohibitions must be closed. In the context of artificial intelligence: all vulnerabilities that could lead to algorithmic harms must be closed.

The rule of verification and confirmation decides that verification and confirmation must be done before making decisions. In the context of artificial intelligence: system safety must be confirmed before deployment, and reliance on assumption is not permitted.

CHAPTER FIVE

REFORMULATION OF JURISPRUDENTIAL RULES IN THE AGE OF AI

This reference presents a reformulation of Islamic jurisprudential rules in the context of proactive jurisprudence for artificial intelligence:

Rule One: If the algorithmic system carries potential catastrophic risks, it must not be deployed until its safety is proven conclusively.

Rule Two: All vulnerabilities that could lead to algorithmic harms must be closed.

Rule Three: Verification and confirmation of algorithmic system safety must be done before deployment, and reliance on assumption is not permitted.

Rule Four: Ethical responsibility falls on developers, operators, and users proportional to their ability to predict risk and prevent harm.

PART TWO

THE PREDICTIVE AND MATHEMATICAL STRUCTURE

VOLUME THREE

THE ALGORITHMIC RISK PREDICTION MODEL

CHAPTER SIX

THE BASIC EQUATION FOR RISK PREDICTION

We present a mathematical model for predicting algorithmic risks before their occurrence.

The Basic Equation:

$$\text{Risk_Score}(a, t) = w_1 \text{ times Complexity}(a) + w_2 \text{ times Autonomy}(a) + w_3 \text{ times Impact_Potential}(a) + w_4 \text{ times Uncertainty}(a) + w_5 \text{ times Historical_Failures}(a)$$

Where: a is the algorithmic system. t is time. $\text{Complexity}(a)$ is the degree of system complexity. $\text{Autonomy}(a)$ is the degree of system independence. $\text{Impact_Potential}(a)$ is the potential impact of harm. $\text{Uncertainty}(a)$ is the degree of uncertainty in system behavior. $\text{Historical_Failures}(a)$ is the number of previous failure cases. w_1 through w_5 are weights determined according to context.

CHAPTER SEVEN DYNAMIC PREDICTION PARAMETERS

We present a dynamic equation because risks change with time, but the current equation is fixed:

$$\text{Risk_Score}(a, t) = \text{Sum of } w_i(t) \text{ times Factor}_i(a, t) \text{ times Decay}(t)$$

Where $\text{Decay}(t)$ is the temporal decay factor that reflects that old data may lose its importance with the passage of time.

CHAPTER EIGHT NETWORK PARAMETERS

We present the systemic risk equation because systems do not operate in isolation but in a connected network:

$$\text{Systemic_Risk}(t) = \text{Sum of Risk_Score}(\text{System}_i, t) \text{ times Connectivity}(i, j)$$

Where $\text{Connectivity}(i, j)$ is the degree of connection between system i and other systems.

CHAPTER NINE EPISTEMIC UNCERTAINTY PARAMETERS

We present the total risk equation because some risks are unknown unknowns:

$\text{Total_Risk} = \text{Known_Risk} + \text{Unknown_Risk} \times \text{Epistemic_Uncertainty}$

Where $\text{Epistemic_Uncertainty}$ measures the degree of our ignorance about unknown risks.

CHAPTER TEN ANOMALY DETECTION ALGORITHMS

We present anomaly detection algorithms in algorithmic system behavior:

$\text{Anomaly_Score}(a, t) = \text{absolute value of Behavior}(a, t) \text{ minus Expected_Behavior}(a) \text{ divided by sigma}$

Where: $\text{Behavior}(a, t)$ is the actual behavior of the system at time t . $\text{Expected_Behavior}(a)$ is the expected behavior. Sigma is the standard deviation.

If Anomaly_Score exceeds a certain threshold, an early warning is issued.

CHAPTER ELEVEN CATASTROPHIC SCENARIOS SIMULATION MODEL

We present a simulation model for potential catastrophic scenarios:

$\text{Catastrophe_Probability}(a) = \text{Sum of Probability}(\text{Scenario_i}) \times \text{Severity}(\text{Scenario_i})$

Where: Scenario_i is the potential catastrophic scenario. $\text{Probability}(\text{Scenario_i})$ is the probability of its occurrence. $\text{Severity}(\text{Scenario_i})$ is its severity.

If $\text{Catastrophe_Probability}$ exceeds a certain threshold, system deployment is prevented.

VOLUME FOUR PROACTIVE SAFETY CERTIFICATES

CHAPTER TWELVE SAFETY CERTIFICATE STANDARDS

This reference presents standards for proactive safety certificates for algorithmic systems:

Standard One: The Complexity Test. Through analyzing the degree of system complexity, evaluating interpretability and explainability, and verifying the existence of black boxes.

Standard Two: The Autonomy Test. Through evaluating the degree of system independence, verifying mechanisms of human oversight, and evaluating the possibility of intervention.

Standard Three: The Impact Test. Through evaluating the potential impact on individuals and society, analyzing catastrophic scenarios, and evaluating the possibility of retreat from decisions.

Standard Four: The Uncertainty Test. Through evaluating the degree of uncertainty in system behavior, analyzing failure cases, and evaluating mechanisms for dealing with potential unknowns.

Standard Five: The History Test. Through analyzing the previous performance record, evaluating previous failure cases, and verifying mechanisms of learning from errors.

CHAPTER THIRTEEN CERTIFICATE GRANTING MECHANISM

Proactive safety certificates are granted through five stages: self-assessment where the developer presents a comprehensive report about the system, simulation testing where the system is tested in a simulation environment, independent evaluation where an independent committee evaluates the system according to standards, final review, and certificate granting if the system meets the standards.

PART THREE THE TECHNOLOGICAL INFRASTRUCTURE

VOLUME FIVE THE DISTRIBUTED ALGORITHMIC ACCOUNTABILITY REGISTRY

CHAPTER FOURTEEN BLOCKCHAIN STRUCTURE

This reference presents a distributed algorithmic accountability registry model based on blockchain technology.

The registry is based on recording every algorithmic decision in distributed blocks on a global network, encrypting records in a way that cannot be modified, providing absolute transparency for all parties, enabling access to records for judges and accredited experts, and with retroactive effect.

Technical specifications include: Permissioned type blockchain, PBFT consensus algorithm, encryption using post-quantum algorithms, geographically distributed nodes, and management by accredited supervisory bodies.

CHAPTER FIFTEEN INTEGRATION WITH MONITORING SYSTEMS

The registry integrates with global monitoring systems: internal monitoring systems for companies, governmental oversight systems, monitoring systems, and public complaint systems.

Contradictions between declared behavior and actual behavior are detected automatically, and penalties are imposed on non-compliant systems.

VOLUME SIX THE ALGORITHMIC EMERGENCY PROTOCOL

CHAPTER SIXTEEN AUTOMATIC ALGORITHMIC CIRCUIT BREAKERS

This reference presents the algorithmic emergency protocol that transforms accountability from reactive to preventive.

The protocol relies on monitoring algorithmic risk indicators in real time, activating automatic algorithmic circuit breakers when indicators exceed a certain threshold, stopping dangerous systems before the occurrence of disaster, and automatically notifying supervisory bodies and courts.

CHAPTER SEVENTEEN MATHEMATICAL FORMULA FOR AUTOMATIC CIRCUIT BREAKERS

Risk severity is calculated as follows:

$$\text{Risk_Index}(t) = \text{Sum of Risk_Score}(\text{System_i}, t) \text{ times Weight}(\text{System_i})$$

Where: System_i is the algorithmic system. Risk_Score(System_i, t) is its risk degree at time t. Weight(System_i) is its weight based on potential impact.

When Risk_Index exceeds a certain threshold, automatic circuit breakers are activated through five steps: immediately stopping dangerous systems, isolating affected systems, preserving

records in the distributed registry, notifying supervisory bodies, and activating compensation mechanisms.

VOLUME SEVEN AI FOR RISK PREDICTION

CHAPTER EIGHTEEN RISK PREDICTION ALGORITHMS

This reference presents the use of artificial intelligence for predicting algorithmic risks through machine learning models to predict system behavior, early warning systems for algorithmic risks, catastrophic scenario prediction algorithms, and simulation models for different scenarios.

CHAPTER NINETEEN DETECTION AND BIAS DISCRIMINATION ALGORITHMS

This reference presents detection and bias discrimination algorithms in algorithmic systems through analyzing training data, detecting biased patterns, evaluating impact on different categories, and issuing automatic alerts for biased systems.

PART FOUR THE LEGAL AND INSTITUTIONAL FRAMEWORK

VOLUME EIGHT THE INTERNATIONAL AI COURT

CHAPTER TWENTY COURT STRUCTURE

This reference presents the establishment of the International Artificial Intelligence Court as a binding legal mechanism.

The court consists of: judges specialized in law and technology, experts in artificial intelligence and ethics, representatives from various geographical regions, and representatives of civil society.

CHAPTER TWENTY-ONE

COURT JURISDICTION

The court is specialized in: adjudicating cross-border algorithmic disputes, applying Elrakhawi Theory equations, issuing legally binding decisions, imposing penalties on non-compliant systems, and overseeing the implementation of compensations.

VOLUME NINE

THE PROACTIVE ETHICS COMMITTEE

CHAPTER TWENTY-TWO COMMITTEE STRUCTURE

This reference presents the establishment of the Proactive Ethics Committee to evaluate the ethical impact of artificial intelligence before its deployment.

The committee consists of: ethics and philosophy scholars, artificial intelligence experts, representatives of civil society, and representatives of affected categories.

CHAPTER TWENTY-THREE COMMITTEE JURISDICTION

The committee is specialized in: evaluating the ethical impact of algorithmic systems before deployment, issuing proactive ethics certificates, proposing long-term ethical policies, monitoring system compliance with standards, and presenting periodic reports on the state of ethics in artificial intelligence.

VOLUME TEN

THE RIGHT TO DIGITAL SAFETY

CHAPTER TWENTY-FOUR LEGAL BASIS FOR THE RIGHT

This reference presents legal recognition of the right to digital safety as a fundamental human right.

This right is based on: the Universal Declaration of Human Rights, the International Covenant on Civil and Political Rights, the principles of precaution and prevention in Islamic jurisprudence, and environmental law.

CHAPTER TWENTY-FIVE RIGHT ENFORCEMENT MECHANISMS

The right to digital safety can be enforced through: national and international judicial lawsuits, complaint mechanisms in the Artificial Intelligence Court, international compensation mechanisms including parametric compensations, and international sanctions on non-compliant systems.

PART FIVE ADDRESSING FUNDAMENTAL GAPS

VOLUME ELEVEN ADDRESSING THE PROBLEM OF FALSE PREDICTION

CHAPTER TWENTY-SIX ACCURATE PREDICTION PARAMETERS

This reference presents the accurate prediction parameters to address the problem of false predictions:

$$PAF(a) = (True_Positives + True_Negatives) / Total_Predictions$$

Where: True_Positives are correct predictions of risk. True_Negatives are correct predictions of safety. Total_Predictions is the total number of predictions.

CHAPTER TWENTY-SEVEN PREDICTION APPEAL MECHANISM

If a developer believes that a false prediction harmed them, they can file an appeal. The appeal is reviewed by an independent prediction committee. If the prediction is proven wrong, the developer is compensated from the prediction insurance fund.

CHAPTER TWENTY-EIGHT PREDICTION INSURANCE FUND

The prediction insurance fund is funded from one percent of safety certificate fees and is used to compensate developers harmed by false predictions.

VOLUME TWELVE ADDRESSING THE PROBLEM OF STIFLED INNOVATION

CHAPTER TWENTY-NINE INNOVATION SUPPORT FUND

This reference presents the innovation support fund to address the problem of stifling innovation:

The fund is funded from two percent of major company profits and provides grants for small and emerging companies to cover safety certificate costs.

CHAPTER THIRTY SIMPLIFIED SAFETY CERTIFICATES

Simplified safety certificates with lower costs and faster procedures are provided for low-risk systems to ensure that the barrier to market entry does not become an obstacle to innovation.

VOLUME THIRTEEN ADDRESSING THE PROBLEM OF ETHICAL DILEMMAS

CHAPTER THIRTY-ONE ETHICAL DILEMMAS PROTOCOL

This reference presents the ethical dilemmas protocol to address the problem of choices between two harms.

The protocol relies on multiple ethics committees including philosophers, legal scholars, religious scholars, and the same experts, public referendums on major dilemmas, and reverse learning algorithms that learn values from human behavior, with the principle of least harm as a default.

CHAPTER THIRTY-TWO REVERSE LEARNING ALGORITHMS

Reverse learning algorithms are used to extract ethical values from actual human behavior instead of imposing programmed values a priori.

VOLUME FOURTEEN ADDRESSING THE PROBLEM OF REPEATED BIAS

CHAPTER THIRTY-THREE ANTI-REPEATED BIAS ALGORITHMS

This reference presents anti-repeated bias algorithms to address the problem of bias in prediction algorithms.

The algorithms rely on: mandatory external audit of training data every six months, reverse feedback mechanisms that automatically correct bias, mandatory diversity algorithms in training data, and algorithmic impartiality certificates.

CHAPTER THIRTY-FOUR ALGORITHMIC IMPARTIALITY CERTIFICATES

Algorithmic impartiality certificates are granted to systems that prove their freedom from bias through repeated and independent tests.

VOLUME FIFTEEN ADDRESSING THE PROBLEM OF TECHNOLOGICAL SOVEREIGNTY

CHAPTER THIRTY-FIVE MULTI-LEVEL GOVERNANCE MODEL

This reference presents the multi-level governance model to address the problem of technological sovereignty.

The model consists of: the national level where each country applies the theory according to its sovereignty, the regional level including the European Union, the African Union, and the Arab League, and the international level which is the International Artificial Intelligence Court. Compliance levels are chosen according to countries' capacities through soft compliance mechanisms.

VOLUME SIXTEEN ADDRESSING THE PROBLEM OF RAPID DEVELOPMENT

CHAPTER THIRTY-SIX DYNAMIC ADAPTATION MECHANISM

This reference presents the dynamic adaptation mechanism to address the problem of rapid development.

The mechanism consists of: periodic review every six months of standards and equations, a continuous update committee of scholars and experts, versioning systems v1.0, v2.0, v3.0, and automatic upgrade mechanisms when new technologies emerge.

VOLUME SEVENTEEN ADDRESSING THE PROBLEM OF PREDICTION RESPONSIBILITY

CHAPTER THIRTY-SEVEN PRINCIPLE OF PREDICTIVE PROPORTIONALITY

This reference presents the principle of predictive proportionality to address the problem of responsibility for predictions.

The principle is based on: no penalty for false predictions unless gross negligence is proven, responsibility is proportional to the degree of certainty, a prediction with ninety percent certainty carries greater responsibility than a prediction with sixty percent certainty, and the prediction insurance mechanism covers predictive risks.

VOLUME EIGHTEEN ADDRESSING THE PROBLEM OF GENERAL AI

CHAPTER THIRTY-EIGHT SUPER GOVERNANCE PROTOCOL

This reference presents the super governance protocol to address the problem of general artificial intelligence.

The protocol relies on: the principle of the human in the loop as a rule that cannot be violated, absolute transparency mechanisms, the super ethics committee that includes philosophers, religious scholars, and lawyers, and mechanisms for stopping emergencies that cannot be disabled for general artificial intelligence, where the AI must reveal every decision it makes.

VOLUME NINETEEN

ADDRESSING THE PROBLEM OF ECONOMIC COST

CHAPTER THIRTY-NINE MIXED FINANCING MODEL

This reference presents the mixed financing model to address the problem of economic cost.

The model consists of: forty percent from rich countries based on historical responsibility, thirty percent from major companies based on profits, twenty percent from taxes on algorithmic transactions, and ten percent from international donations and grants.

VOLUME TWENTY ADDRESSING THE PROBLEM OF PRIVACY AND TRANSPARENCY

CHAPTER FORTY GRADUATED TRANSPARENCY MODEL

This reference presents the graduated transparency model to address the problem of privacy and transparency.

The model consists of: the first general level that includes basic information about the system, the second restricted level that includes detailed information accessible only to accredited experts, and the third secret level that includes trade secrets that are not disclosed except by judicial order, with mechanisms for knowledge verification that confirm technical accuracy without disclosing details.

PART SIX PRACTICAL APPLICATIONS

VOLUME TWENTY-ONE REAL CASE STUDIES

CHAPTER FORTY-ONE CASE STUDY: THE FLASH CRASH 2010

On the sixth of May 2010, the Dow Jones index lost approximately one trillion dollars in minutes due to the interaction of dozens of automated trading algorithms.

Applying the Elrakhawi Theory: it was possible to predict risks before the disaster, it was possible to activate automatic circuit breakers, it was possible to record every algorithmic decision in the distributed registry, and it was possible to prevent the disaster before its occurrence.

CHAPTER FORTY-TWO CASE STUDY: THE UBER INCIDENT 2018

In March 2018, a self-driving car developed by Uber killed a woman in Arizona.

Applying the Elrakhawi Theory: it was possible to predict system risks before deployment, it was possible to refuse granting the safety certificate, it was possible to activate automatic circuit breakers upon risk detection, and it was possible to record every decision for accountability in the distributed registry.

CHAPTER FORTY-THREE CASE STUDY: AMAZON ALGORITHM BIAS

In 2018, Amazon discovered that its recruitment algorithm was biased against women.

Applying the Elrakhawi Theory: it was possible to detect bias before deployment through detection algorithms, it was possible to prevent harm before its occurrence, it was possible to automatically compensate those harmed, and it was possible to refuse granting the ethics certificate.

CHAPTER FORTY-FOUR CASE STUDY: GENERATIVE AI RISKS

In 2024, generative artificial intelligence models produced misleading images and texts that affected elections.

Applying the Elrakhawi Theory: it was possible to predict misinformation risks before deployment, it was possible to activate automatic verification mechanisms, it was possible to hold developers accountable through the distributed registry, and it was possible to prevent the spread of misleading information.

VOLUME TWENTY-TWO SECTOR DIVISION FOR APPLICATIONS

CHAPTER FORTY-FIVE FINANCIAL SERVICES

The Elrakhawi Theory is applied to the financial sector through: proactive safety certificates for trading algorithms, a distributed registry to track every trading decision, automatic circuit breakers to prevent disasters, and parametric compensations for victims.

CHAPTER FORTY-SIX HEALTHCARE

The Elrakhawi Theory is applied to the health sector through: proactive safety certificates for diagnostic algorithms, verification mechanisms of medical recommendations with ninety-nine percent accuracy, a distributed registry to track every medical decision, and parametric compensations for medical errors.

CHAPTER FORTY-SEVEN SMART TRANSPORT

The Elrakhawi Theory is applied to smart transport through: proactive safety certificates for self-driving cars, a distributed registry to track every driving decision, automatic circuit breakers to prevent accidents, and parametric compensations for victims.

CHAPTER FORTY-EIGHT SMART RECRUITMENT

The Elrakhawi Theory is applied to smart recruitment through: proactive ethics certificates for sorting algorithms, bias detection algorithms, a distributed registry to track every recruitment decision, and parametric compensations for those harmed.

PART SEVEN CRITICISM AND CONSTRAINTS

VOLUME TWENTY-THREE SCIENTIFIC CHALLENGES

CHAPTER FORTY-NINE PREDICTION CHALLENGES

The theory faces challenges in accurate risk prediction: uncertainty in complex system behavior, difficulty predicting unprecedented scenarios, complex interactions between different factors, and the need for more historical data.

The proposed solution includes: developing more advanced models, using deep learning techniques, collecting more data, and adding epistemic uncertainty parameters to the research.

VOLUME TWENTY-FOUR POLITICAL CHALLENGES

CHAPTER FIFTY RESISTANCE OF MAJOR COMPANIES

The theory faces resistance from major companies: fear of regulatory constraints, lack of desire for transparency, commercial sovereignty, lobbying, and pressure from interest groups.

The proposed solution includes: gradual implementation, economic incentives for compliant companies, international popular pressure, and cooperation with pioneering companies.

CHAPTER FIFTY-ONE GEOPOLITICAL CHALLENGES

The theory faces geopolitical challenges: tensions between major countries, lack of balance in international power, difficulty of international consensus, and regional conflicts.

The proposed solution includes: building international alliances, engaging all parties, focusing on common interests, and using mediation mechanisms.

CONCLUSION

The Elrakhawi Theory of Proactive Jurisprudence for Artificial Intelligence presents a deep and methodological solution to a fundamental problematic in the age of artificial intelligence: How do we build strong and responsible accountability systems before the occurrence of crises?

The theory is based on eighteen innovative principles including proactive jurisprudence, the algorithmic risk prediction model, proactive safety certificates, the distributed algorithmic accountability registry, the algorithmic emergency protocol, the proactive ethics committee, the principle of digital precaution, the right to digital safety, the accurate prediction parameters, the innovation support fund, the ethical dilemmas protocol, anti-repeated bias algorithms, the multi-level governance model, the dynamic adaptation mechanism, the principle of predictive

proportionality, the super governance protocol, the mixed financing model, and the graduated transparency model.

The theory presents a complete technological framework including an advanced mathematical model for risk prediction that includes dynamic prediction parameters, network parameters, and epistemic uncertainty parameters, along with anomaly detection algorithms, the cognitive robot, the distributed accountability registry, the proactive ethics committee, emergency protocols, and proactive safety certificates systems.

The theory presents a complete legal framework including the right to digital safety, enforcement mechanisms, and international sanctions.

The theory integrates with the European Artificial Intelligence Act and emerging laws in China and the United States by presenting a methodological solution to the problematic of reactive accountability through the shift to predictive governance.

The theory relies on Islamic jurisprudence rich in the rules of precaution, prevention, blocking the means, and verification, presenting a civilizational vision for accountability in the age of artificial intelligence.

It is not merely an abstract theory. The Elrakhawi Theory constitutes a fundamental building block in constructing a safer and more just future in the age of artificial intelligence. It is a practical tool, not just an academic one, that can be used immediately by countries, companies, and international organizations to build strong proactive accountability systems.

APPENDICES

APPENDIX ONE: COMPLETE MATHEMATICAL DETAILS FOR MODELS

Risk Prediction Equation: $\text{Risk_Score}(a, t) = w_1 \text{ times Complexity}(a) + w_2 \text{ times Autonomy}(a) + w_3 \text{ times Impact_Potential}(a) + w_4 \text{ times Uncertainty}(a) + w_5 \text{ times Historical_Failures}(a)$

Dynamic Prediction Equation: $\text{Risk_Score}(a, t) = \text{Sum of } w_i(t) \text{ times Factor}_i(a, t) \text{ times Decay}(t)$

Systemic Risk Equation: $\text{Systemic_Risk}(t) = \text{Sum of Risk_Score}(\text{System}_i, t) \text{ times Connectivity}(i, j)$

Total Risk Equation: $\text{Total_Risk} = \text{Known_Risk} + \text{Unknown_Risk times Epistemic_Uncertainty}$

Anomaly Detection Equation: $\text{Anomaly_Score}(a, t) = \text{absolute value of Behavior}(a, t) \text{ minus Expected_Behavior}(a) \text{ divided by sigma}$

Catastrophe Simulation Equation: $\text{Catastrophe_Probability}(a) = \text{Sum of Probability}(\text{Scenario_i}) \times \text{Severity}(\text{Scenario_i})$

Risk Index Equation: $\text{Risk_Index}(t) = \text{Sum of Risk_Score}(\text{System_i}, t) \times \text{Weight}(\text{System_i})$

Accurate Prediction Parameters: $\text{PAF}(a) = (\text{True_Positives} + \text{True_Negatives}) / \text{Total_Predictions}$

APPENDIX TWO: BIAS DETECTION ALGORITHMS

The algorithms work as follows: collecting training data, analyzing demographic distribution, detecting biased patterns, evaluating impact on different categories, issuing automatic alerts for biased systems, and imposing penalties on biased systems.

APPENDIX THREE: PROACTIVE SAFETY CERTIFICATES STRUCTURE

The certificate includes five basic components: the five assessment criteria, the independent and self-assessment mechanism, simulation testing, and final review and granting.

APPENDIX FOUR: THE DISTRIBUTED ALGORITHMIC ACCOUNTABILITY REGISTRY

The technical structure includes: Permissioned type blockchain, PBFT consensus algorithm, encryption using post-quantum algorithms, geographically distributed nodes.

Mathematical formula for decision recording: $\text{Decision_Hash} = \text{SHA3}(\text{algorithm_id} + \text{decision_data} + \text{timestamp} + \text{context})$

APPENDIX FIVE: DETAILED CASE STUDIES

Case Study One: The Flash Crash 2010. Case Study Two: The Uber Incident 2018. Case Study Three: Amazon Algorithm Bias. Case Study Four: Generative AI Risks 2024. Case Study Five: Medical Error Due to Diagnostic Algorithm. Case Study Six: Deep Forgery and Impact on Elections.

APPENDIX SIX: INTEGRATION WITH ISLAMIC JURISPRUDENCE

The rule of repelling corruption takes precedence over bringing benefit and its application. The rule of blocking the means and its application. The rule of verification and confirmation and its

application. The rule of harm is not removed by harm and its application to algorithmic accountability.

APPENDIX SEVEN: ROADMAP TO GLOBAL IMPLEMENTATION

Phase One (1-2 years): Academic and technical establishment.

Phase Two (2-5 years): Experimental application.

Phase Three (5-10 years): Global application.

Phase Four (10+ years): Civilizational consolidation.

REFERENCES

LEGAL REFERENCES:

1. Donoghue v Stevenson, AC 562, House of Lords, 1932.
2. Palsgraf v Long Island Railroad Co, 248 NY 339, 1928.
3. European Parliament. 2024. Regulation EU 2024/1689 on Artificial Intelligence, AI Act.
4. United Nations. 2023. Global Compact on Artificial Intelligence.

ARTIFICIAL INTELLIGENCE REFERENCES:

5. Russell, S., and Norvig, P. 2020. Artificial Intelligence: A Modern Approach. Fourth edition. Pearson.
6. Goodfellow, I., Bengio, Y., and Courville, A. 2016. Deep Learning. MIT Press.
7. Bostrom, N. 2014. Superintelligence: Paths, Dangers, Strategies. Oxford University Press.
8. Floridi, L., et al. 2018. AI4People: An Ethical Framework for a Good AI Society. Minds and Machines 28(6): 1203-1221.

AI ETHICS REFERENCES:

9. Jobin, A., Ienca, M., and Ravizza, P. 2019. The Global Landscape of AI Ethics Guidelines. Nature Machine Intelligence 1(9): 389-399.
10. Mittelstadt, B., et al. 2016. The Ethics of Algorithms: Mapping the Debate. Big Data and Society 3(2).

COMPLEX SYSTEMS REFERENCES:

11. Barabasi, A. L. 2016. Network Science. Cambridge University Press.

12. Newman, M. 2010. Networks: An Introduction. Oxford University Press.

GAME THEORY REFERENCES:

13. Osborne, M. J., and Rubinstein, A. 1994. A Course in Game Theory. MIT Press.

14. Myerson, R. B. 1991. Game Theory: Analysis of Conflict. Harvard University Press.

ISLAMIC JURISPRUDENCE REFERENCES:

15. Ibn Qudamah. Al-Mughni. Dar al-Kitab al-Arabi.

16. Al-Nawawi. Majmu Sharh al-Madhhab. Dar al-Fikr.

17. Ibn Taymiyyah. Majmu al-Fatawa. Majma Malik Fahd.

18. Al-Zuhayli. Al-Fiqh al-Islami wa Adillatuhu. Dar al-Fikr.

19. Al-Shatibi. Al-Muwafaqat. Dar Ibn Affan.

REAL CASE REFERENCES:

20. Securities and Exchange Commission. 2010. Findings Regarding the Market Events of May 6, 2010.

21. National Transportation Safety Board. 2019. Collision Between Vehicle Controlled by Developmental Automated Driving System and Uber Test Vehicle.

22. Dastin, J. 2018. Amazon Scraps Secret AI Recruiting Tool that Showed Bias Against Women. Reuters.

BLOCKCHAIN REFERENCES:

23. Nakamoto, S. 2008. Bitcoin: A Peer-to-Peer Electronic Cash System. White Paper.

24. Buterin, V. 2014. Ethereum: A Next-Generation Smart Contract Platform. White Paper.

25. Castro, M., and Liskov, B. 1999. Practical Byzantine Fault Tolerance. OSDI.

INDEX

A: Digital Safety, Proactive Ethics, Digital Precaution, Anomalous Patterns

B: Blockchain, Algorithmic Emergency Protocol

C: Algorithmic Bias, Risk Prediction

D: Algorithmic Trust

E: Proactive Jurisprudence

F: Case Studies, Right to Digital Safety

G: Detection Algorithms, Automatic Circuit Breakers

H: Repelling Corruption, Distributed Accountability Registry
I: Generative Artificial Intelligence
J: Reactive Accountability, Proactive Accountability
K: Blocking the Means
L: Proactive Safety Certificates
M: Algorithmic Transparency
N: Parametric Compensation Fund
O: Algorithmic Quality Assurance
P: Prediction Methods
Q: Smart Contracts, International AI Court
R: Accountability Gap
S: Proactive Jurisprudence, Islamic Jurisprudence
T: Precaution Rules, Repelling Rules
U: Flash Crash
V: Centralization
W: Proactive Ethics Committee, Risk Prediction Equation, Digital Precaution Principle
X: Elrakhawi Theory, Predictive Models
Y: Legal Engineering
Z: Developer Obligations
AA: Algorithmic Certainty

COPYRIGHT AND INTELLECTUAL PROPERTY NOTICE

All intellectual property rights are fully reserved to the author, Dr. Mohamed Kamal Arafa Elrakhawi.

Academic and scientific quotation is permitted for research and criticism purposes, provided that the source is cited completely and clearly in accordance with academic standards, mentioning the author name, reference name, and digital identifier.

Academic citation format: Elrakhawi, Mohamed Kamal. 2026. The Elrakhawi Theory of Proactive Jurisprudence for Artificial Intelligence: From Reactive Accountability to Predictive Governance in the Age of Complex Intelligent Systems. DOI: 10.5281/zenodo.21868244.

Any reproduction, production, manufacture, commercial exploitation, or reverse engineering of any part of this reference, including mathematical equations, algorithms, legal models, case studies, and texts, is strictly prohibited without obtaining explicit written consent directly from the author.

Any industrial or commercial use that relies on equations, algorithms, models, or texts contained in this reference will expose its owner to international legal accountability under intellectual property protection laws and author rights.

This reference has been deposited and temporally registered as a definitive proof of scientific priority for the author.

DOI: 10.5281/zenodo.21868244

Date of Publication: July 2026

Author: Dr. Mohamed Kamal Arafa Elrakhawi

All Rights Reserved

END OF DOCUMENT