

Algorithmic Decision-Making in the Digital Age: Legal Accountability, Transparency and Human Rights

Dr. Lukas Schneider¹, Dr. Marek Kowalski²

¹ Institute for Technology Law and Digital Governance, Berlin, Germany.

² Institute of Cyber Law and Emerging Technologies, Warsaw, Poland.

Received: 18 Jan 2024 Revised: 28 Jan 2024 Accepted: 28 Jan 2024 Published: 29 Jan 2024

ABSTRACT

The rapid integration of artificial intelligence and algorithmic systems into governmental, commercial, financial, employment, and social environments has transformed the manner in which decisions are made. Automated systems can process vast quantities of information and generate decisions more rapidly than traditional human processes. However, algorithmic decision-making also creates significant legal concerns relating to transparency, accountability, discrimination, privacy, explainability, and human oversight. This article examines the legal implications of automated and semi-automated decision-making, focusing on situations in which algorithmic outputs may significantly affect individuals' rights and interests. It analyzes the challenges associated with opaque algorithms, biased datasets, automated profiling, insufficient explanations, and difficulties in assigning responsibility when decisions are produced by complex technological systems. The article argues that technological efficiency should not override fundamental legal principles of fairness, accountability, and human dignity. It proposes a governance framework based on transparency, risk assessment, human oversight, explainability, data governance, independent auditing, documentation, and effective mechanisms for challenging automated decisions. The article concludes that algorithmic systems should operate within a rights-oriented legal framework in which technological innovation is accompanied by meaningful accountability and accessible human review.

Keywords: Artificial Intelligence; Algorithmic Decision-Making; Technology Law; Human Rights; Algorithmic Accountability; Transparency; Data Governance.

© 2024. The Author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution 4.0 International License (CC BY 4.0), which permits unrestricted use, distribution and reproduction in any medium, provided the original work is properly cited.

1. Introduction

Artificial intelligence has moved from experimental technological environments into everyday decision-making systems.

Algorithms are increasingly used by financial institutions, employers, public authorities, online platforms, healthcare organizations, educational institutions, and other entities.

These systems can analyze enormous quantities of information and identify patterns that may be difficult for humans to recognize.

For organizations, algorithmic decision-making can improve efficiency, reduce administrative costs, and support faster decisions.

However, automated decision-making creates an important legal problem.

When an algorithm makes a decision affecting an individual, it may be difficult to determine **why the decision was made, who is responsible for it, and how the affected individual can challenge it.**

A conventional human decision-maker can generally provide an explanation for a decision.

An algorithmic system may instead produce an output based on complex statistical relationships that are difficult to interpret.

This creates the possibility of an accountability gap.

The central issue is therefore not whether algorithms should be used, but how their use can be reconciled with legal principles of fairness, transparency, equality, privacy, and human dignity.

This article examines these challenges and proposes a framework for accountable algorithmic decision-making.

2. Concept of Algorithmic Decision-Making

Algorithmic decision-making occurs when computer-based systems process information and generate an outcome, recommendation, classification, or decision.

A simplified model is:

Data → Algorithmic Processing → Risk/Prediction → Decision → Outcome

The degree of human involvement may vary.

Some systems merely provide recommendations to human decision-makers.

Others automatically determine outcomes without direct human intervention.

This distinction is legally important because the potential impact on individuals increases as human involvement decreases.

3. Automated and Semi-Automated Decisions

Algorithmic systems can generally be divided into two broad categories.

Automated Decision-Making

The system independently produces and implements the decision.

Semi-Automated Decision-Making

The system generates a recommendation, while a human makes the final decision.

Semi-automated systems may appear to preserve human control.

However, human decision-makers may sometimes rely heavily on algorithmic recommendations.

Therefore, meaningful human involvement requires more than simply placing a person at the end of an automated process.

4. Algorithmic Transparency

Transparency is one of the central challenges associated with AI systems.

An individual affected by an algorithmic decision may reasonably ask:

- What information was used?
- What factors influenced the decision?
- Why was the particular outcome generated?

- Who is responsible for the system?
- Can the decision be challenged?

Organizations should therefore maintain sufficient documentation concerning the operation and purpose of important algorithmic systems.

Transparency does not necessarily require disclosure of confidential source code.

Instead, individuals should receive meaningful information about how the system affects them.

5. The Problem of the Black Box

Some AI systems are described as "black boxes" because their internal decision-making processes may be difficult to understand.

Deep-learning systems can contain large numbers of computational parameters.

Even system developers may find it difficult to provide a simple explanation for an individual prediction.

This creates a tension between technological complexity and legal accountability.

A decision affecting an individual's rights should not become unreviewable merely because it was produced through sophisticated technology.

6. Algorithmic Bias

Algorithmic bias represents another major concern.

AI systems learn from data.

If training data contains historical inequalities or incomplete representation, an algorithm may reproduce or amplify those patterns.

Bias can potentially affect:

- recruitment;
- financial services;
- insurance;
- education;
- policing;
- public services; and
- online content systems.

The use of an algorithm does not automatically make a decision objective.

Technological systems can reproduce human biases present in their underlying data.

7. Data Quality

The reliability of algorithmic decisions depends significantly upon data quality.

Poor-quality data may result in:

- inaccurate predictions;
- incorrect classifications;
- outdated information;
- duplicate records; or
- unfair outcomes.

Organizations should therefore implement data-quality procedures before deploying high-impact AI systems.

The principle can be summarized as:

Poor Data → Poor Model Output → Potentially Unfair Decision

8. Privacy and Profiling

Algorithmic systems frequently rely on large quantities of personal information.

They may create profiles based on:

- purchasing behavior;
- online activity;
- location;
- financial information;
- professional history; or
- interaction patterns.

Profiling can provide legitimate benefits, but excessive profiling may interfere with individual privacy.

Organizations should therefore carefully evaluate the necessity and proportionality of data processing.

9. Automated Decisions in Employment

Employers increasingly use technology during recruitment and workforce management.

Algorithms may assist in:

- screening applications;
- identifying candidates;
- evaluating performance;
- predicting employee turnover; and
- allocating work.

However, employment decisions directly affect individuals' economic opportunities.

An algorithmic system that systematically disadvantages particular groups could create serious legal consequences.

Human review and regular bias testing should therefore be considered essential for high-impact employment applications.

10. Algorithmic Credit Decisions

Financial institutions may use algorithms to assess creditworthiness.

Automated systems can analyze large quantities of financial information and generate rapid decisions.

However, an individual may be denied financial services without understanding the underlying reason.

An effective framework should provide mechanisms for:

- reviewing relevant information;
 - correcting inaccurate data;
 - obtaining meaningful explanations; and
 - challenging potentially incorrect decisions.
-

11. Government Use of Algorithms

Public authorities may use algorithms for administrative decision-making.

Potential applications include:

- identifying potential fraud;
- allocating public resources;
- processing applications;
- risk assessment; and
- detecting irregularities.

Government use requires particularly strong safeguards because individuals may have limited ability to choose whether to interact with the relevant system.

Public-sector algorithmic systems should therefore be subject to strong documentation, oversight, and accountability mechanisms.

12. Human Oversight

Human oversight should be meaningful rather than symbolic.

A human reviewer should have sufficient authority and information to:

- question the algorithmic recommendation;
- examine relevant evidence;
- identify potential errors;
- reject the algorithmic output; and
- provide an alternative decision.

If the human simply accepts every algorithmic recommendation, the process may remain effectively automated.

13. Explainability

Explainability concerns the ability to communicate why an algorithm produced a particular outcome.

Different levels of explanation may be appropriate depending on the circumstances.

A basic explanation might identify the major factors influencing the decision.

A technical explanation may describe the model architecture and variables.

The appropriate level should depend upon:

- the significance of the decision;
- the complexity of the system;
- the rights affected; and
- the potential consequences of error.

14. Right to Challenge

An algorithmic decision should not necessarily be final.

Individuals affected by significant automated decisions should have appropriate mechanisms to challenge them.

A possible process is:

Automated Decision → Notification → Explanation → Human Review → Correction/Confirmation → Appeal

Such mechanisms can reduce the consequences of algorithmic errors.

15. Accountability

A central legal question is identifying who is responsible for an algorithmic decision.

Potentially relevant actors include:

- AI developers;
- technology providers;
- system operators;
- data controllers;
- organizations deploying the system; and
- human decision-makers.

Responsibility should be clearly allocated before an AI system is deployed.

Organizations should not be able to avoid responsibility simply by arguing that "the algorithm made the decision."

16. Algorithmic Auditing

Independent algorithmic audits can identify problems before they cause significant harm.

Audits may examine:

- accuracy;
- discrimination;
- data quality;
- security;
- explainability;
- performance;
- documentation; and
- compliance.

High-risk systems should undergo periodic assessment rather than being evaluated only once before deployment.

17. Security of Algorithmic Systems

AI systems can themselves become targets for cyberattacks.

Attackers may attempt to manipulate:

- training data;
- model inputs;
- system parameters;
- access credentials; or
- output processes.

Security therefore forms an essential component of algorithmic governance.

A system that is accurate under normal conditions may become unreliable if its inputs are deliberately manipulated.

18. Documentation and Recordkeeping

Organizations should maintain appropriate documentation throughout the AI lifecycle.

Records may include:

- purpose of the system;
- training-data sources;
- model version;

- testing results;
- risk assessments;
- human oversight procedures;
- system updates; and
- significant decisions.

Documentation can be valuable when investigating disputes or regulatory concerns.

19. Proposed Governance Framework

A responsible algorithmic decision-making framework can be organized into eight stages.

Stage 1 – Purpose Assessment

Identify why the AI system is necessary.

Stage 2 – Risk Classification

Determine the potential impact of algorithmic decisions.

Stage 3 – Data Assessment

Evaluate the quality, relevance, and reliability of data.

Stage 4 – Model Testing

Test accuracy, robustness, and potential bias.

Stage 5 – Human Oversight

Establish meaningful human review mechanisms.

Stage 6 – Transparency

Provide appropriate information concerning the system's role.

Stage 7 – Continuous Auditing

Monitor performance after deployment.

Stage 8 – Remedy

Provide accessible mechanisms for correction and appeal.

20. Risk-Based Regulation

Not every algorithm requires the same level of legal oversight.

A recommendation system suggesting movies may create relatively limited consequences.

A system determining access to employment, credit, essential services, or significant public benefits can have substantially greater impact.

Regulatory requirements should therefore be proportional to risk.

A possible classification is:

Risk Level	Example	Suggested Oversight
Low	Entertainment recommendation	Basic transparency
Medium	Marketing profiling	Privacy and monitoring
High	Recruitment screening	Human review and auditing
Very High	Significant public or legal decisions	Strong human oversight and independent review

21. Privacy by Design

Privacy should be integrated into algorithmic systems from the beginning.

A privacy-oriented system should consider:

- data minimization;
- access controls;
- encryption;
- limited retention;
- secure processing; and
- appropriate anonymization or pseudonymization.

Privacy should not be treated as an afterthought.

22. Future of Algorithmic Governance

Algorithmic decision-making is likely to expand into additional areas.

AI systems may increasingly assist with:

- legal analysis;
- healthcare decisions;
- financial assessment;
- public administration;
- education;
- employment; and
- security.

As these applications expand, the distinction between technological assistance and autonomous decision-making will become increasingly important.

Future governance frameworks will need to evolve accordingly.

23. Conclusion

Algorithmic decision-making offers significant technological advantages.

AI systems can process information rapidly, identify patterns, support complex analysis, and improve organizational efficiency.

However, algorithmic systems can also create significant legal risks when their decisions affect individuals' rights and opportunities.

The principal challenges include opacity, bias, inaccurate data, privacy concerns, inadequate human oversight, cybersecurity vulnerabilities, and uncertainty concerning responsibility.

The solution is not to prohibit algorithmic decision-making altogether.

Instead, organizations should adopt a rights-oriented governance framework based on transparency, risk assessment, data quality, explainability, human oversight, independent auditing, security, documentation, and effective mechanisms for challenging decisions.

The principle of **meaningful human accountability** should remain central.

Technology may generate recommendations or even automate certain decisions, but organizations deploying such systems must remain capable of explaining, reviewing, and, where necessary, correcting their outcomes.

Ultimately, responsible algorithmic governance requires a balance between **innovation and accountability**. Algorithms can improve decision-making, but their legitimacy depends upon ensuring that efficiency never becomes a justification for unfairness, discrimination, or the removal of meaningful human control.

Declarations

Conflict of interest: The authors declare that they have no conflict of interest.

Funding: This review received no specific grant from any funding agency in the public, commercial or not-for-profit sectors.

Ethical approval: Not applicable. This article is a review of published literature and did not involve human participants or animal subjects.

Data availability: No new data were generated or analyzed in support of this review. All sources relied upon are cited in the References.

Use of AI tools: No generative artificial intelligence tool was used in the conception, drafting or analysis of this article. The authors take full responsibility for the content.

Author contributions: Author contributed to the conception of the review, the survey of the literature, and the drafting and revision of the manuscript. Author read and approved the final version.