

Legal Issues and Regulation of AI

Ameen Jauhar

Artificial Intelligence – An Overview

- Lack of uniform or universally accepted definition of what technologies constitute AI.
- Instead certain facets and functionalities can be used to identify AI
- *Mimicking human cognition or features of human cognition, through data-intensive algorithms, trained on large data corpuses to perform specific tasks.*
- Narrow AI (aka task-specific AI) and General AI (AGI).

Application of Narrow AI

Majority of current AI systems are deployed for one or more of the following functions:

- Generative AI (like ChatGPT/GPT-4, Dall-E);
- Predictive algorithms (like COMPAS used in US parole systems);
- Profiling (like credit scoring algorithms used for determining credit worthiness of potential consumers); and
- Autonomous Decision-Making (like recruitment algorithms targeted in the proposed NY legislation, or content moderation algorithms on social media platforms like Twitter).

Legal Issues in AI

- Fundamental rights based risks
 - Violation of FRs guaranteed in Part III of the Constitution
 - Example – biases which violate A. 14
 - Uses that can stifle free speech (surveillance, predictive policing tools like FRT)
 - Privacy concerns
 - Large datasets scraped/accessed for processing
 - Profiling
 - Data breaches
 - Absence of proportionality in certain applications

Legal Issues (contd.)

- Procedural Issues
 - Specific uses, especially where legal rights of a person can be affected, must be guided through a statutory framework.
 - Example – use of algorithms to profile and authenticate identity of pensioners. Such a use has obvious privacy considerations since it collects and processes sensitive personal data (biometrics), but also has the potential to impact the legal right of pensioner.
 - Legislation must set out clear guardrails to mitigate risks of bias, inaccurate outcomes (false negatives), and prescribe clear redressal framework.

Overview of AI regulation in India

- NITI Aayog Responsible AI principles
- Indications of general algorithmic regulation principals in the reported Digital India Act
- Intermediary Guidelines 2021 incidentally brought in the concept of human-in-loop (through human review) of content moderation.
- State policy/principle frameworks (like in Telangana – TAIM).
- Overall absence of any coherent sectoral, or overarching regulatory framework.
- Additional non binding guidelines, toolkits, frameworks (NITI on FRT, NASSCOM on industry, ICMR on biomedical research).

The European Union Artificial Intelligence Act

- Applies to all EU Member States, AI systems in the EU Market and AI systems affecting the EU population
- Governance to be performed by national market surveillance authorities
- A European Artificial Intelligence Board proposed for implementation and development of standards
- Regulatory approach is the imposition of Civil Liability
- Objectives:
 - AI Systems should be safe and legally compliant
 - Facilitate Innovation and Investment in AI
 - Enhance Governance and Enforcement
 - Develop Single Market and Prevent Market Fragmentation

The EU AI Act's Risk Based Approach

- All AI are classified into four tiers: Unacceptable Risk, High Risk, Limited Risk and Minimal Risk.
- High Risk is AI used in critical sectors (law enforcement, infrastructure, justice etc.)
 - Obligation on High-Risk AI to conduct risk assessment; use high quality datasets; and ensure traceability, compliance, human oversight and high levels of robustness and security.
- Limited Risk AI : transparency obligation – informing the user that they are conversing with an AI to enable informed decision making
- Minimal Risk AI : No regulatory interference with free use AI like spam filters or AI-enabled video games.
- Enforcement of fines is also classified into tiers
 - Unacceptable risk fined at 30 million euros or 6% Global Annual Turnover
 - Other breaches at 20 million euros or 4% Global Annual Turnover
 - Incorrect, Incomplete, Misleading Information to the competent authority at 10 million euros or 2% Global Annual Turnover

The Chinese Algorithm Framework

- China's Cyberspace Administration (CAC) has released draft guidelines (w.e.f. March 2022) to regulate algorithmic recommender systems.
- These regulations demand new technical requirements and auditing of keywords for greater user access and control.
- These regulations also demand adherence to 'mainstream values' and preventing the use of algorithms for illicit behaviour.

Recommender Algorithm Systems are used by social media, e-commerce, rideshare, food delivery and targeted advertising companies

Illicit Behaviour may refer to anti-competitive practices or price discrimination based on personal characteristics

'Mainstream Values' refer to using AI for the public good, reducing consumption and over-indulgence in the populace.

The Chinese Algorithm Framework

- Enforcement includes fines, debarring offenders from enrolling new users, having licenses revoked and having websites/apps shut down.
- The rules apply to internet information service providers – a term which covers any organization providing online services.
- The rule's objectives cite national security, standardization of algorithm regulations and preventing harm to individuals.
- The CAC has also sought to regulate the use of deepfakes – to prevent disruption of social welfare.

The USA's AI Initiative Act 2020

- Current US regulation lacks a comprehensive AI law and regulates AI is through the cross-application of rules such as product liability, data privacy, intellectual property, discrimination, and workplace rights
- The FTC regulation mandates 'truth, fairness and equity' by asking AI to do 'more good than harm'.
- Over the past two years, more than 30 bills referencing AI have been introduced in state legislatures.
- The Algorithmic Accountability Act has recently been reintroduced into the US senate. Senators Klobuchar and Lummis have also introduced a bill to combat social media algorithms.
- Also, the OSTP published an AI Bill of Rights in 2022, bringing a human-centric approach to AI development and deployment.
- New York introduced Local Law 144 of 2021 to prohibit ADMs in recruitment processes. However, yet to enforce it.

The Brazilian Approach

- The Brazilian congress recently passed a bill that develops an AI framework in the nation. It must now pass the senate.
- The framework is inspired by the recommendations on AI by the OECD (human-centric development of AI).
- This framework is also supported by the Brazilian Strategy for AI, a broad policy guideline set earlier.
- The AI framework requires transparency and informing users about the institution responsible for developing the AI System.
- Places an emphasis on human rights, democratic values and environmental preservation
- It proposes a risk based management approach and specifies sector-based intervention.
- It has been criticised for its generality and lack of technical grounding.

The Australian Approach

- Developed a voluntary, 8-step AI Ethics Framework.
- The framework calls for mitigating safety risks of unreliable AI, a multi-step transparency process and contestability among others.
- Relies on the OECD framework as well as the EU's ethics guidelines for trustworthy AI.
- The Australian Human Rights Commission has made AI liability recommendations through remedies in already existent legislations.
- Materiality threshold for regulatory action – AI must have a legal or similarly significant effect on an individual.
- Recommends strong regulation on FRT and biometric technologies, and a moratorium on these technologies till such a moratorium is in place.

Singaporean Approach

- No intention of adopting robust regulation, instead choosing to go with a prescriptive, voluntary framework.
- Model Governance Framework released in 2019 calls for AI to be 'human centric' and 'explainable, transparent and fair'.
 - The 4 pillars of the framework are internal governance structures, determining the level of human involvement in AI decision making, operations management and stakeholder interaction.
 - Risk based management approach is recommended

UK's AI White Paper

- Cross-sectoral approach
- Focus on promoting incentives for innovation to facilitate commercial domination of Britain in the AI industry
- Limited focus on risk-mitigation (no substantive inputs on what regulatory frameworks must look like)
- Left to independent regulators to formulate actual instruments and frameworks of governance.
- Fails to offer anything significant (beyond principles) in terms of regulation.

Contrasting Approaches

Jurisdiction	The EU	China	USA	Brazil	Singapore	Australia	UK
Objective	Economic Incentives	Societal Welfare	Emphasis on non-discrimination	Focus on ethical usage (environmental impact)	No discernible unique focus	Focus on preventing state abuse	Promoting innovation and responsible adoption
Regulatory Approach	Mandatory , robust framework with civil penalties	Mandatory framework with deterrent penalties	Lacks a central framework that dictates a regulatory approach	General framework that lacks technical grounding	Self Governance based model	Voluntary Framework	Sectoral regulation
Management Style	Risk Based Approach	Managemen t of AI based on moral and societal impact	‘More good than harm’ – based on relative harm	Risk based management	Risk Based Approach	Risk Based Approach	Strong impetus towards the commercial exploitation of AI

Thank you

