



Regulating AI-powered consumer contracts: risks, remedies, and global legal responses

Regolamentazione dei contratti dei consumatori basati sull'Intelligenza Artificiale: rischi, rimedi e risposte giuridiche globali

ETTORE MARIA LOMBARDI 

Associate Professor of Private Law and Digital Society Law
University of Florence

Abstract

Artificial Intelligence (AI) has reshaped consumer contracts by enabling automated drafting, personalization, and enforcement at scale, with limited human intervention. While this increases efficiency, it also introduces risks, including bias, opacity, privacy violations, and power asymmetries. This article offers a comparative legal analysis of regulatory responses and case studies, focusing on the EU AI Act and contrasting it with approaches in the US, Canada, Brazil, and emerging frameworks in Asia-Pacific. The paper also highlights best practices and regulatory interventions, from mandatory human oversight to third-party audits, to support a harmonized global governance.

L'Intelligenza Artificiale (IA) ha trasformato i contratti dei consumatori, consentendo la redazione, la personalizzazione e l'enforcement su larga scala, con un limitato intervento umano. Sebbene tali sviluppi accrescano l'efficienza, ad essi conseguono rilevanti criticità, tra cui bias algoritmici, opacità decisionale, violazioni della privacy e asimmetrie di potere. Il contributo propone una analisi giuridica comparata delle risposte regolatorie, anche attraverso case study, con particolare riferimento all'AI Act dell'Unione Europea, posto a confronto con gli approcci adottati negli Stati Uniti, in Canada, in Brasile e nelle legislazioni emergenti dell'area Asia-Pacifico. L'articolo evidenzia, inoltre, best practice e strumenti di intervento regolatorio, tra cui la supervisione umana obbligatoria e gli audit indipendenti, valorizzando una governance globale armonizzata.



Keywords: EU AI Act; GDPR; Data Protection; Contract Law; Consumer Law

Summary: [1. Introduction: The Role of AI in the Day-By-Day Consumer Transactions.](#) – [2. The Rise of AI Consumer Contracting.](#) – [3. Data Protection and Privacy Breaches.](#) – [4. A Comparative Regulatory Analysis: the European Union AI Act and GDPR Interplay, the United States' Sectoral Enforcement and Fragmentation, Canada's AIDA and a Hybrid Oversight, the Brazilian Civil Liability and High-Risk Prohibitions, the Asia-Paci](#) – [5. A Common Trend: Diverging Tools, Converging Themes.](#) – [6. Liability, Remedies, and Enforcement.](#) – [7. Case Law and Enforcement Trends in the European Union and in the U.S.A.: Enforcement Under the AI Act and GDPR, on the one side, FTC and CFPB Enforcement on the other.](#) – [8. United Kingdom and The Algorithmic Insurance Discrimination.](#) – [9. A Hypothesis of Multi-Jurisdiction Litigation: Dynamic Pricing Cause and Emerging Enforcement Patterns.](#) – [10. Toward an emerging Global Convergence?](#) – [11. Integrating AI Regulation into Consumer and Contract Law Reform.](#) – [12. Toward Responsible AI in Consumer Contracts: Synthesis and Future Outlook.](#)

1. Introduction: The Role of AI in the Day-By-Day Consumer Transactions.

Artificial intelligence has swiftly migrated from the margins of legal operations into the heart of everyday consumer transactions. Enabled by large language models (LLMs), deep learning algorithms, and autonomous agents, AI systems now routinely manage the entire lifecycle of contracts, from drafting to execution, without direct human involvement.

These developments promise significant efficiency gains and mass customization of terms, but they also raise urgent legal and ethical questions.¹ Can consumers truly consent to contracts generated by opaque machine processes? Are existing doctrines of unconscionability, fairness, and transparency equipped to govern algorithmic decision-making? And how should liability be allocated when an AI-generated clause causes consumer harm?

Nowhere are these concerns more acute than in the context of consumer contracting since these are high-volume, low-negotiation transactions where bargaining power already skews heavily in favor of corporations.² The introduction of AI threatens to widen this gap, enabling firms to deploy tailored but inscrutable contractual terms that elude regulatory scrutiny and judicial review.

Moreover, empirical studies have shown that a substantial portion of consumers cannot meaningfully comprehend or challenge AI-generated

¹ See G Alpa, *What Is Private Law?* (Pacini Giuridica 2010) passim, but especially 9 ff, where the A. analyses the functions of private law and the need to recalibrate doctrines when social and technological conditions change.

² See, for further details, G Alpa, *Responsabilità dell'impresa e tutela del consumatore* (Giuffrè 1975) passim, where the A. analyses enterprise liability and consumer protection in standard contracts.

clauses, and this profile raises doubts about the validity of consent and the enforceability of these agreements under traditional principles of contract law. For instance, a 2023 survey across various EU states found that less than 43% of respondents reported that they feel in full control of the content shown and the decisions that they make online.³ Such findings amplify concerns that “consent by confusion” may undermine the meeting of minds essential to contract formation.⁴

This article addresses these challenges through a comprehensive legal analysis of AI-powered consumer contracts. It explores how different jurisdictions have responded, identifies regulatory gaps and convergence trends, and offers doctrinal and policy-based recommendations for building a trustworthy and legally sound ecosystem of automated contracting. Special attention is given to the European Union’s AI Act, which categorizes AI-based contract tools as “high-risk” and imposes transparency, accountability, and human oversight obligations.

This analysis is contextualized through comparative references to the U.S., Canada, Brazil, and select Asia-Pacific jurisdictions, allowing for a nuanced understanding of global governance trajectories.⁵

2. The Rise of AI Consumer Contracting.

The digitization of legal processes has entered a new phase with the deployment of AI systems capable of autonomously managing high-volume contract workflows. In sectors such as e-commerce, insurance, and personal finance, algorithms now perform tasks historically handled by legal professionals: drafting standard form agreements, customizing clauses, monitoring compliance, and even renegotiating renewal terms. For example, e-commerce platforms, such as Amazon and insurance providers in the UK, deploy contract-generation bots that operate in real time, tailoring contract terms based on user interaction data and risk assessments derived from

³ BEUC - The European Consumer Organisation, *Connected, but Unfairly Treated: Consumer Survey Results on the Fairness of the Online Environment* (Brussels 2023) 8. See, further, C Scognamiglio, ‘Condizioni generali di contratto nei rapporti tra imprenditori e la tutela del “contraente debole”’ (1987) 2 Riv dir comm 425 ff, where the A. analyses the control of general terms and weaker-party protection, a doctrinal bridge toward algorithmically generated “take-it-or-leave-it” terms.

⁴ See G Perlingieri, *Appunti sul contratto telematico* (Edizioni Scientifiche Italiane 2000), that represents an early work addressing how digital contracting challenges the traditional “meeting of minds” in contract formation; G Conte, ‘La formazione del contratto’ in FD Businelli (ed), *Comm. Schlesinger* (Giuffrè 2018) passim, where the A. examines classical consent requirements in contract formation.

⁵ L Balestra, ‘An Introduction to the Importance of a Comparative Approach in the Analysis of the Benefits and the Risks of E-commerce’ in G Finocchiaro, L Balestra and M Timoteo (eds), *Major Legal Trends in the Digital Economy: The Approach of the EU, the US, and China* (Il Mulino 2022) 23 ff, where the A. highlights the importance of comparative analysis in digital contract regulation; M Andenas and M Heidemann (eds), *Quo Vadis Commercial Contract? Reflections on Sustainability, Ethics and Technology in the Emerging Law and Practice of Global Commerce* (Springer 2023), where it is emphasized a comparative and innovative perspectives on adapting contract law to new technologies like AI.

machine learning models.⁶ These AI-driven systems allow businesses to issue personalized contracts to millions of consumers efficiently, but they also reduce opportunities for human review before terms are agreed, reproducing a sort of automation at scale.⁷

At the heart of this transformation are generative models, particularly LLMs, which can process complex legal inputs and produce human-like contractual text. LLMs such as GPT-5 or Claude have been integrated into contract lifecycle management tools, enabling dynamic contract creation with limited legal oversight. These systems not only draft clauses but also interpret terms, propose revisions, and even simulate negotiation outcomes based on historical data.

Even more advanced are agentic AI systems – autonomous software agents programmed to perform legal tasks iteratively, with an evolving understanding of counterparty behavior and regulatory constraints.⁸ Such AI agents can, for instance, adjust contract terms on the fly (e.g. interest rates or service levels) in response to user behavior or external data (like market indices), all without human sign-off on each change.⁹

The COVID-19 pandemic catalyzed the mainstream adoption of AI in digital transactions, because businesses rapidly transitioned to remote-first operations, relying on AI-driven contracting tools to replace manual review and paper-based workflows.¹⁰ However, this accelerated deployment has outpaced legal adaptation, leaving regulators scrambling to retrofit old doctrines to new technologies.¹¹ Already some years ago, indeed, it was noticed that European contract law's traditional paradigms would be tested by these autonomous processes, presaging the extensive reforms now underway.¹² In short, the post-2020 era witnessed an unprecedented scaling-up of automated contracting that represents a trend that is likely irreversible and demands urgent legal scrutiny regarding the legal risks in AI-Powered consumer contracts.

In this sense, transparency deficits and explainability challenges raise strongly, especially if the hallmark of contract law is considered to be the

⁶ G Sartor, 'Artificial Intelligence and Human Rights: Between Law and Ethics' (2020) 27(6) Maastricht J Eur & Comp L 705 ff.

⁷ See, further, M Maggiolo, *Il contratto predisposto* (Cedam 1996), which represents a seminal monograph on pre-formulated standard contracts and their legal treatment.

⁸ L Edwards and M Veale, 'Slave to the Algorithm? Why a "Right to an Explanation" Is Probably Not the Remedy You Are Looking For' (2017) 16 Duke L & Tech Rev 27.

⁹ See G Gitti, 'L'applicazione dei sistemi di intelligenza artificiale nei contratti per l'impresa' (2023) 2 Tecnologie e diritto 339 ff, where the A. analyses the AI systems in enterprise contracting and the resulting reallocations of responsibility and information duties; G Gitti, 'Robotic Transactional Decisions' Osservatorio del diritto civile e commerciale 619 ff, where the A. treats the decision automation in transactions and develops a work conceptually relevant for 'agentic' AI negotiating and executing consumer contracts.

¹⁰ From mortgage agreements to healthcare consent forms, AI-generated contracts became a structural necessity for continuity.

¹¹ F Capriglione, 'Fintech: un laboratorio per i giuristi' (2019) 2 Contratto e impresa 377 ff, where the A. notes that rapid fintech innovation reveals the limits of traditional regulatory models; M Ebers, 'Liability for Artificial Intelligence and EU Consumer Law' JIPITEC 204 ff.

¹² Ebers, 'Liability for Artificial Intelligence' (n 4) 208 ff; P Perlingieri, 'Presentazione' in P Perlingieri, S Giova and I Prisco (eds), *Atti del XV Convegno SISDiC – "Rapporti civilistici e intelligenze artificiali: attività e responsabilità"* (2020) VII ff, where the A. observes that AI-driven processes test classical contract-law paradigms.

mutual assent, or a meeting of the minds. Yet in AI-powered consumer contracting, this principle is increasingly undermined by the opacity of algorithmic systems, due to the fact that many AI tools, especially those based on neural networks or deep learning, operate as “black boxes” that produce outputs without intelligible explanations of their decision-making logic.¹³ Consequently, critical terms in consumer contracts may be shaped by variables the consumer (or even the business deploying the AI) cannot easily perceive or understand, and consequently this lack of transparency violates duties of clarity and openness embedded in consumer protection law and general contract doctrine. Under EU law, for example, sellers are required to draft terms in plain, intelligible language and ensure that consumers are informed of the terms’ implications. An AI system that dynamically generates or alters terms based on undisclosed factors runs afoul of these transparency obligations.¹⁴

This issue becomes even more pressing when AI is used to personalize pricing, risk scoring, or eligibility conditions, considering, for example, as subscription agreements with dynamic pricing clauses generated by AI may adjust rates based on undisclosed consumer behavior patterns, without any meaningful disclosure to the consumer. Such practices clash with fundamental transparency requirements and potentially breach laws like Article 5(1)(a) of the GDPR, which mandates fairness and transparency in personal data processing, and Article 6(1) of the EU Unfair Contract Terms Directive, which mandates that unfair terms not bind consumers.¹⁵ From a contract law perspective, if consumers are not aware that an algorithm is dictating key terms or on what basis it is doing so, genuine informed consent is questionable, and traditional defense, such as claiming that the consumer “agreed” by clicking a box, lose force when the substantive terms are crafted by an opaque AI rather than presented in a stable form open to negotiation or review.¹⁶

The issues described above raise problems related to algorithmic bias and discriminatory outcomes, especially because AI systems trained on historical data often replicate and exacerbate societal biases. This phenomenon is

¹³ S Wachter, B Mittelstadt and C Russell, ‘Counterfactual Explanations Without Opening the Black Box: Automated Decisions and the GDPR’ (2018) 31 Harv J L & Tech 841 ff.

¹⁴ European Commission, ‘Guidance on the Interpretation and Application of Directive 93/13/EEC on Unfair Terms in Consumer Contracts’ OJ C323/4 [https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:52019XC0927\(01\)](https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:52019XC0927(01)) accessed 30 June 2026. See, for further details, R Di Raimo, ‘Finanza derivata e consenso contrattuale: osservazioni a valle delle crisi d’inizio millennio’ *Dialoghi di Diritto dell’Economia* 1 ff, where the A. highlights that in contracts marked by information asymmetry and no negotiation, true consent is often supplanted by formalistic safeguards.

¹⁵ See Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation) OJ L119/1, art 5(1)(a); Council Directive 93/13/EEC of 5 April 1993 on unfair terms in consumer contracts OJ L95/29, art 6(1).

¹⁶ See M Maggiolo, *Servizi ed attività di investimento: prestatori e prestazione* (Giuffrè, 2012), for the contractual architecture of investment services and suitability/information duties, relevant where AI systems shape consumer financial contracts; G Gitti, ‘L’applicazione dei sistemi di intelligenza artificiale nei contratti per l’impresa’ (2023) 2 *Tecnologie e diritto*, 339 ff., where the A. underscores the need to reconcile AI-driven contract personalization with transparency and consent requirements.

acutely visible in the contracting context. For instance, in 2023 the U.S. Consumer Financial Protection Bureau (CFPB) has held lenders accountable for discriminatory credit scoring.¹⁷ In the U.K., consumer advocates and regulators uncovered that certain postcode-based exclusions in auto insurance contracts (determined by algorithmic models) disproportionately penalized ethnic minority communities in urban centers.¹⁸ These outcomes arise when AI proxies, alike ZIP codes or online behavior data, correlate with protected characteristics, resulting in discriminatory impacts even without explicit intent.

Such biased contractual outcomes are not merely ethical lapses as they contravene anti-discrimination laws and may violate fundamental rights. In the EU, for example, Article 21 of the Charter of Fundamental Rights prohibits discrimination on grounds such as race or ethnicity, a principle that extends to algorithmic decision-making by virtue of overarching equality laws.

The EU's AI Act accordingly classifies AI systems used in important areas of opportunity (credit, employment, housing, etc.) as "high-risk" precisely because of the potential for such disparate impacts, and these systems will require pre-deployment bias testing, ongoing monitoring, and human oversight to ensure compliance.¹⁹ Likewise, U.S. regulators have signaled that algorithmic decisions in credit or insurance must comply with existing civil rights statutes, like the Equal Credit Opportunity Act, enacted October 28, 1974, that makes it unlawful for any creditor to discriminate against any applicant, with respect to any aspect of a credit transaction, on the basis of race, color, religion, national origin, sex, marital status, or age (provided the applicant has) in lending and state insurance anti-discrimination laws, regardless of the novelty of the technology. In short, an AI-generated contract term that, say, charges higher prices to residents of certain neighborhoods or profiles users by race is likely to be legally unenforceable and expose its deployer to liability.²⁰

¹⁷ Consumer Financial Protection Bureau, 'CFPB acts to protect the public from black-box credit models' (2022) <https://www.consumerfinance.gov/about-us/newsroom/cfpb-acts-to-protect-the-public-from-black-box-credit-models-using-complex-algorithms/> accessed 30 June 2026. See, for further details, DS Cromvello, 'Algorithmic Discrimination and Equal Protection Law: Legal Remedies for AI Bias in Automated Decision-Making' (2025) 1 Int J Law Policy & Sci Res 6 ff.,.

¹⁸ Financial Conduct Authority (UK), 'Potential Bias in Firms' Use of Consumer Data (Research Report)' (September 2023) <https://www.fca.org.uk/panels/consumer-panel/potential-bias-firms-use-consumer-data> accessed 30 June 2026, where it is noted evidence of an "ethnicity penalty" in insurance pricing.

¹⁹ C Perlingieri, 'Intelligenza Artificiale tra principi e regole' (2025) 1 Int Rev Artif Intell Law 35 ff, where the A. emphasizes the EU's attempt to balance fundamental rights and innovation in its AI regulatory framework.

²⁰ P Perlingieri, 'Constitutional Norms and Civil Law Relationships' Ital J 1 ff, where the A. develops further the thesis of the "constitutionalisation" of private law, relevant for reading AI-contract disputes through dignity, equality and substantive fairness; P Perlingieri, *Il diritto civile nella legalità costituzionale* (Edizioni Scientifiche Italiane 2020) passim, where the A. argues that private-law interpretation must be consistent with constitutional principles, especially where vulnerable parties and fundamental rights are implicated; G Alpa, *The Right to Be Oneself* (Hart Publishing 2024), where the A. develops a rights-based perspective on identity and dignity, resonant with algorithmic discrimination and 'personalisation' in contracting.

3. Data Protection and Privacy Breaches, Consent and Comprehension Failures.

AI contract systems rely heavily on personal data, often extending beyond what is necessary for lawful contract formation, and data minimization principles under Article 5(1)(c) GDPR are routinely flouted when LLMs or agentic AIs ingest extensive behavioral, biometric, or financial datasets for “personalization” purposes without a clear legal basis. In 2023 alone, several major European financial institutions reported GDPR-reportable incidents where contract-generation bots inadvertently exposed or misused sensitive personal data.²¹

Such breaches raise complex questions of liability and oversight, especially since, under the GDPR, the deploying company remains the data controller responsible for how personal data is used in contract drafting, and, if the AI’s operation leads to unlawful processing (e.g., using irrelevant or excessive personal data in setting terms), regulatory penalties and damages claims can follow. Moreover, the EU AI Act introduces an additional layer of accountability as providers of high-risk AI systems will be obliged to ensure robust data governance (including quality, relevance, and minimization of training data) and could face strict liability for damages caused by non-compliance in data management.²² Notably, Article 71 of the AI Act foresees a liability regime whereby the provider of a non-compliant high-risk AI system can be held liable for harm caused, without the need for the injured party to prove negligence.²³

This, combined with existing privacy laws, means AI contracting systems operate under a dense web of data protection duties, and consequently a failure in any link – be it a privacy breach or an AI compliance gap – can render the resulting contract void or unenforceable and trigger substantial fines.

Another critical risk concerns the erosion of meaningful consent since AI-generated contracts are frequently written in dense “legalese” or generated in real-time, making it difficult for the average consumer to comprehend the obligations being imposed. Studies have documented that consumers already struggle with standard form contracts, and the complexity multiplies when terms can vary from person to person or moment to moment due to algorithmic personalization.²⁴ If an AI system produces a bespoke set of terms for each user, perhaps varying arbitration clauses or renewal conditions based on predicted customer behavior, the traditional notion that a consumer had a fair chance to review and agree to the terms becomes tenuous. Further, many such

²¹ For example, a bank’s AI tool that auto-filled agreement templates drew on an internal customer profile database and accidentally embedded credit score data and health information into a loan contract, exceeding what the customer had consented to share. See European Data Protection Board, *Annual Report 2023: Safeguarding Individuals’ Digital Rights* (Bruxelles 2023) 24 ff.

²² See, on this topic, C Scognamiglio, ‘Responsabilità civile ed intelligenza artificiale: quali soluzioni per quali problemi?’ *Resp civ e prev* 1073 ff, where the A. maps liability models for AI risk, including fault, strict liability, and evidentiary presumptions.

²³ See, further, *ibid* 1073 ff, where the A. analyzes emerging liability models for AI-driven harms.

²⁴ BEUC, *Artificial Intelligence: What Consumers Say. Findings and Policy Recommendations of a Multi-country Survey on AI* (3 September 2020) 6 https://www.beuc.eu/sites/default/files/publications/beuc-x-2020-078_artificial_intelligence_what_consumers_say_report.pdf accessed 30 June 2026.

contracts are presented via clickwrap or browsewrap interfaces²⁵ where hurried users give perfunctory consent, unaware of unique terms crafted for them by an AI.²⁶

Although contract law has long wrestled with the problem of informed consent in mass contracting, such as the battle against unread end-user license agreements, for example, AI threatens to deepen that problem: when each consumer's contract is slightly different and generated by an algorithm, consumers cannot rely on community knowledge or external summaries of "the standard terms," because there may be none. This raises the specter of what might be called "algorithmic adhesion contracts," where the adhesion is not just because of take-it-or-leave-it format, but because the terms are dynamically opaque, and, in legal terms, these practices risk invalidation. In this regard it is relevant to underline that, under Article 3(1) of the EU Consumer Rights Directive, consent must be given via a clear affirmative act after being informed of the contractual terms. Similarly, the Unfair Terms Directive allows terms to be struck out if they were not transparently communicated or if they create a significant imbalance to the detriment of the consumer. If AI complexity is weaponized to secure nominal consent, courts may treat such consent as vitiated by misinformation or misunderstanding, and, as a consequence, one could argue that an AI-drafted term which a consumer could not realistically understand might be deemed unenforceable as an "unreadable" or non-transparent term, much as hidden fees or surprising clauses have been under existing law.

4. A Comparative Regulatory Analysis: the European Union AI Act and GDPR Interplay, the United States' Sectoral Enforcement and Fragmentation, Canada's AIDA and a Hybrid Oversight, the Brazilian Civil Liability and High-Risk Prohibitions, the Asia-Pacific as a patchwork in progress.

The EU Artificial Intelligence Act, formally adopted in June 2024, represents the first comprehensive legislative framework governing AI systems. Central to its architecture is a risk-based approach: AI systems are classified into tiers (unacceptable risk, high-risk, limited risk, and minimal risk) with "high-risk" systems subject to the most stringent requirements, and AI systems used in consumer contracting, particularly those that determine significant eligibility, pricing, or terms for access to essential private services (financial services, utilities, etc.), are explicitly deemed high-risk under Annex III of the AI Act.²⁷ For example, an AI system that evaluates consumers for creditworthiness or

²⁵M Garcia, 'Browsewrap: A unique solution to the slippery slope of the clickwrap conundrum' (2013) 36 Campbell L Rev 31 ff.

²⁶RJ Girasa, 'Click-Wrap, Shrink-Wrap, and Browse-Wrap Agreements: Judicial Collision with Consumer Expectations' (2002) 10 North East J Legal Stud 102.

²⁷*Regulation (EU) 2024/1689 (Artificial Intelligence Act)*, Annex III (high-risk use-cases), point 6, where AI systems for assessing consumer creditworthiness, insurance risk, etc., are classified as high-risk.

insurance coverage, or that automatically sets contractual terms in those domains, falls within the high-risk category.²⁸

More in particular, high-risk AI systems must comply with an array of ex ante obligations, including, i) pre-deployment conformity assessments to ensure the system meets EU standards before it can be placed on the market (AI Act Article 19), ii) data governance and management requirements (Articles 10 and 11), such as using training data that is relevant, representative, and free of errors or bias, and observing data minimization in line with the GDPR, iii) transparency and explainability (Article 13), meaning the system should be designed to permit the interpretation of its outputs and to provide users with information about its functioning, iv) human oversight mechanisms (Article 14), ensuring that appropriate human intervention can monitor and override AI decisions that are erroneous or harmful.

Notably, the AI Act complements the GDPR rather than replaces it since the latter continues to govern the processing of personal data within AI systems (e.g. requiring a legal basis for processing, honoring data subject rights, etc.), while the AI Act overlays sector- and use-case-specific rules aimed at the AI's reliability and accountability. Notably an important intersection is with GDPR's Article 22, which gives individuals the right not to be subject to decisions based solely on automated processing that significantly affect them, unless certain safeguards are in place.²⁹ The AI Act bolsters this by requiring, in Article 52, that users of high-risk AI inform affected persons that AI is involved in generating decisions or content, which includes contract terms, and by reinforcing the availability of human review for such decisions. Thus, a consumer presented with an AI-generated contract in the EU would have both the GDPR-based right to receive an explanation or human intervention upon request, and the AI Act-based assurance that the system went through compliance checks for transparency and fairness.

The enforcement regime under the AI Act is robust, reflecting the high stakes of AI failures, since national market surveillance authorities will oversee compliance, with the European Commission coordinating for consistency. Penalties for non-compliance scale with severity as certain violations, like using banned AI practices or failing to meet high-risk requirements, can trigger administrative fines up to 30 million euros or 7% of global annual turnover (whichever is higher)³⁰, surpassing even the GDPR's fine levels for the most

²⁸ See M Heidemann, *Methodology of Uniform Contract Law: The UNIDROIT Principles in International Legal Doctrine and Practice* (Springer 2007), where the A. develops a methodological toolkit for uniform contract interpretation, which is useful especially when AI-driven contracting raises cross-border coordination problems; M Andenas and G Deipenbrock (eds), *Regulating and Supervising European Financial Markets: More Risks than Achievements* (Springer 2016), for a critical perspective on risk governance and supervisory architecture, illuminating how the EU's AI governance model fits broader regulatory design; Finocchiaro, Balestra and Timoteo (n 5), where it is developed a comparative background on platform governance, data and digital markets, contextualising divergent tools and converging themes.

²⁹ See, on this topic, F Capriglione, *Manuale di diritto bancario e finanziario* (Cedam 2024), where the A. develops a comprehensive account of financial-market regulation, including the growing role of automated processes and FinTech, which is useful especially when AI sets credit/insurance terms.

³⁰ *Artificial Intelligence Act*, Article 71(3)(b), which provides fines up to 30 million euros or 7% of annual worldwide turnover for certain serious infringements.

egregious breaches. This signals that the EU intends to treat AI governance on par with, if not more stringently than, data protection enforcement, and, indeed, a company deploying an AI-driven contract platform in Europe must now navigate a dual compliance challenge, both data protection and AI regulation, and faces significant financial and reputational risks for failures in either realm.

The United States lacks a unified AI law comparable to the EU's AI Act, and instead, AI in consumer contracts is addressed (to the extent it is) through a patchwork of sector-specific laws and agency enforcement. In this regard, the divergence between the EU and U.S. approaches to AI governance is not merely regulatory, but also institutional. The EU framework remains anchored in a preventive and *ex ante* compliance-oriented logic, pursued through centralized and horizontal regulatory standards applicable across Member States. By contrast, the U.S. model is built with a predominantly reactive and enforcement-driven approach relying on federal agencies to develop guidelines and enforcement standards while promoting, at the same time, self-regulation within the market. Accordingly, AI regulation is based on a decentralized and sector-specific structure, under which federal agencies have jurisdiction over specific sectors.³¹

Particularly three federal regulators form a triad of oversight in practice composed by the Federal Trade Commission (FTC),³² the Consumer Financial Protection Bureau (CFPB),³³ and the Department of Justice (DOJ).³⁴

Enforcement in the U.S. is largely post hoc, with regulatory interventions typically triggered by consumer complaints, whistleblowers, or ex post investigations rather than pre-market approvals.³⁵

Legislative efforts at the federal level remain nascent, and bills such as the Algorithmic Accountability Act and the Transparent Automated Governance (TAG) Act have been introduced in Congress, which would require companies to conduct impact assessments and disclose their use of automated decision systems that significantly affect consumers. However, as of 2025 these bills

³¹ T Davtyan, 'The U.S. Approach to AI Regulation: Federal Law, Policies, and Strategies Explained' (2025) 16 (2) J Law Tech & Internet, 226 ff.

³² The FTC polices unfair or deceptive acts and practices in commerce, has increasingly scrutinized AI-driven business practices under its broad mandate to protect consumers. This includes false or misleading claims about AI services and the use of AI in ways that harm consumers (e.g. algorithmic fraud or dark pattern marketing). See US Federal Trade Commission, 'Bringing Dark Patterns to Light' (Staff Report 2022) https://www.ftc.gov/system/files/ftc_gov/pdf/P214800+Dark+Patterns+Report+9.14.2022+-+FINAL.pdf accessed 23 April 2026.

³³ The CFPB regulates consumer financial products and services, monitors AI use in credit scoring, loan underwriting, and similar financial contract contexts for compliance with fair lending laws (such as the Equal Credit Opportunity Act) and consumer protection standards.

³⁴ Particularly its *Civil Rights Division*, which can bring actions against companies for discrimination in housing, lending, or employment when AI tools result in disparate impacts on protected groups.

³⁵ For instance, in August 2025, the FTC fined two companies of a total of 145 million US dollars to settle that they misled millions of consumers seeking to buy comprehensive health insurance. The FTC alleged that both companies deceived consumers and led them to purchase plans that did not provide the promised health care coverage, and bombarded consumers with telemarketing and robocalls. See *Federal Trade Commission v MediaAlpha Inc* 2:25-cv-07263 (US Central District of California).

have not passed, leaving businesses to navigate guidance rather than binding law.³⁶ In the meantime, some state legislatures have enacted targeted laws, such as, for example, Illinois' Artificial Intelligence Video Interview Act for hiring tools, or state-level privacy laws that touch on automated profiling, but these are piecemeal. As already noticed, the net result is fragmentation, and a company using AI for consumer contracts in the U.S. must consult a matrix of potentially applicable rules, from FTC Act §5 (unfair/deceptive practices) to sectoral regulations like the Fair Credit Reporting Act (if credit data is involved) or housing discrimination laws, without a single unifying statute or regulator. This fragmentation also extends to remedies since an affected consumer might lodge a complaint with the CFPB, or bring a private lawsuit under anti-discrimination law, or rely on the FTC to act, depending on the issue.

The U.S. remedial enforcement model, while offering regulatory flexibility, presents challenges for consumer protection compared to the preventive and compliance-oriented EU legal framework. Consequently, in the absence of consolidated legislation, effective AI governance will depend on coordination among the federal government, state regulators and industry.

Canada has adopted a hybrid model through its proposed Artificial Intelligence and Data Act (AIDA), which seeks to integrate AI governance with data protection principles. AIDA (tabled in 2022 as part of Bill C-27) introduces a framework somewhat akin to the EU approach, categorizing AI systems by risk and imposing requirements on high-impact systems. Although not yet in force as of 2026, AIDA foreshadows obligations including i) transparency since organizations deploying high-impact AI systems would need to notify individuals that they are interacting with an AI system and provide a general explanation of how it makes decisions, ii) human review rights since individuals significantly affected by an automated decision would have the right to request human review (this aligns with principles found in GDPR Article 22), iii) impact assessments since before deploying a high-impact AI system, companies would have to conduct algorithmic impact assessments evaluating potential harms, bias, and privacy implications, and mitigate identified risks, iv) oversight by the Privacy Commissioner since Canada's privacy regulator is poised to gain authority to audit AI systems and ensure compliance, building on its existing role under privacy laws.

Notably, even before AIDA's enactment, Canada has implemented the Directive on Automated Decision-Making for federal government agencies, which classifies automated systems by impact level and requires algorithmic impact assessments (AIAs) for high-impact systems.³⁷ This directive, while binding only on government uses, has set a de facto benchmark that influences private sector best practices. For example, banks and insurers in Canada are

³⁶ Algorithmic Accountability Act of 2022, HR 6580, 117th Cong (2022) (US House bill proposing mandatory impact assessments for automated decision systems); Transparent Automated Governance Act of 2023, S 1865, 118th Cong (2023) (US Senate bill proposing transparency for AI use in government services).

³⁷ Treasury Board of Canada Secretariat, 'Directive on Automated Decision-Making' (2019, amended 2022), mandating Algorithmic Impact Assessments for automated systems used by federal government.

beginning to voluntarily conduct AIAs for new high-risk AI tools to anticipate future regulatory expectations.

Canada's approach thus mixes hard law with soft guidance. The Office of the Privacy Commissioner has actively opined on AI issues, indicating that existing laws (like the Personal Information Protection and Electronic Documents Act, PIPEDA) already impose constraints on AI use of personal data, and that companies should not wait for AIDA to implement fairness and accountability measures. Enforcement in Canada, pending AIDA, can occur under general consumer protection or human rights statutes. For instance, if an AI underwriting tool were found to discriminate, a provincial human rights commission could take action under anti-discrimination law. Meanwhile, competition law has been floated as a tool against algorithmic collusion or overly dominant AI platforms. In summary, Canada is moving toward a comprehensive AI regulatory framework that, once fully in place, will create compliance parallels with the EU (transparency, oversight, assessment) but in the interim relies on a mosaic of directives and general laws.

Brazil's approach to AI regulation has been notably pro-consumer and precautionary, and indeed, in 2023, Brazil's Congress was considering an AI Bill (Projeto de Lei nº 2338/2023)³⁸ that takes a strict stance on high-risk AI in consumer context³⁹ and whose key features include i) explicit prohibitions on certain AI practices,⁴⁰ ii) strict civil liability for AI providers and users,⁴¹ iii) regulatory authority.⁴²

Even ahead of a dedicated AI law, Brazilian consumer protection law (itself quite strong, via the Consumer Defense Code) and civil rights law provide avenues to tackle AI-driven harms,⁴³ and Brazilian courts have shown willingness to apply broad constitutional principles, such as human dignity and equality, in private law disputes, which could extend to algorithmic discrimination cases. Moreover, collective redress is a hallmark of Brazilian consumer law, and entities like the Public Prosecutor's Office or consumer associations can bring class actions on behalf of consumers. This means that if

³⁸ See below.

³⁹ Projeto de Lei nº 2338/2023 (Brazilian Artificial Intelligence Bill) arts 9–12 (banning certain high-risk AI practices) and arts 19–22 (strict liability provisions for AI-related harm).

⁴⁰ For example, any AI system that exploits a consumer's vulnerable condition (due to age, disability, or illiteracy) to materially distort their decision-making in a contract could be outright banned. Discriminatory algorithms in credit or insurance would similarly be disallowed. Brazil's bill is influenced by fundamental rights concerns and seeks to draw red lines around unacceptable uses of AI.

⁴¹ The draft law would impose liability on those who develop or deploy AI systems that cause harm, regardless of fault in some cases. In particular, if an AI system generates a contract term deemed abusive or causes financial or dignity harm to consumers (through bias or error), the consumer could claim damages without needing to prove negligence—similar to a product liability regime. This approach is underpinned by the idea that AI risks should be borne by those best positioned to prevent them (the companies), not by consumers.

⁴² The bill contemplates a national AI authority to oversee implementation, with powers to investigate, require audits of AI systems, and issue fines. Given Brazil's experience with a proactive data protection authority under its LGPD (General Data Protection Law), one might expect a similarly empowered AI regulator if the bill becomes law.

⁴³ For example, if an insurance contract written by AI contains an opaque clause that harms consumers, that clause could be nullified under the Civil Code's general clauses on good faith and fair balance, or under specific consumer law provisions against abusive terms.

an AI system inflicts a standard harm on many consumers, say, an unfair fee hidden by an algorithm in e-commerce contracts, a collective lawsuit could address it. The prospect of such collective accountability reinforces the need for companies to tread carefully with AI deployment in Brazil.

In the Asia-Pacific region, regulatory maturity on AI and consumer protection varies widely, and Singapore, for example, has championed an agile governance approach. Instead of heavy-handed laws, it issued voluntary Model AI Governance Frameworks (first released in 2019, updated subsequently) which encourage organizations to implement accountability, transparency, and human-centric design in AI systems. Singapore's Personal Data Protection Commission has also published specific guidance on AI, and the government is considering formal regulation, but for now the emphasis is on industry adherence to best practices. In consumer contracts, Singapore relies on its existing Consumer Protection (Fair Trading) Act to tackle unfair practices, and, in this light, an AI clause that is egregiously unfair or misleading could be challenged under that Act. Additionally, Singapore's courts have not yet faced major AI-contract disputes, but the legal infrastructure is in place via contract law and data protection law, which includes obligations for AI-driven decision-making affecting individuals.

On the other hand, Japan has adopted AI governance guidelines (the AI Social Principles and related guides) that promote fairness, transparency, and accountability, but these remain non-binding. Sector-specific regulators in Japan, such as the Financial Services Agency, have nudged industries to ensure AI algorithms, like those used in insurance underwriting or credit scoring, do not violate existing laws or ethical norms. For instance, the FSA has highlighted that insurance AI pricing must still comply with the Insurance Business Act's requirements of fairness and explanation to policyholders. As of 2025, Japan has not passed a comprehensive AI law, and the approach is to adapt existing legal concepts (e.g., the Civil Code for contracts, the APPI for data privacy) to AI contexts case by case.

China has moved quickly to regulate specific facets of AI, particularly algorithms that affect large numbers of consumers, and in this regard the Cyberspace Administration of China issued the Provisions on the Management of Algorithmic Recommendation Systems (effective March 2022), which impose requirements on companies that use personalization algorithms in consumer-facing services.⁴⁴ These rules require companies to disclose basic principles of their algorithms, allow users to opt-out of personalized recommendations, and prohibit algorithmic practices that endanger national security or public order (broadly defined). For consumer contracts, this means Chinese e-commerce or financial platforms using AI to set terms or prices must abide by transparency rules and avoid prohibited behavior, such as, for example, using algorithms to show different prices to different consumers unfairly is under scrutiny. Furthermore, China's Personal Information Protection Law (2021) has provisions on automated decision-making, providing that individuals have the right to not be subjected to price discrimination through automated decision-

⁴⁴ Cyberspace Administration of China, Provisions on the Management of Algorithmic Recommendation Systems (effective 1 March 2022), requiring transparency and non-discrimination in consumer-facing recommendation algorithms.

making and can request explanations. Enforcement in China has already seen regulators penalize companies for algorithmic misdeeds, and, for instance, the government has cracked down on lending apps that used opaque algorithms leading to usurious terms for vulnerable borrowers.

Australia to date has relied on its consumer law and privacy law, with regulators like the Australian Competition and Consumer Commission (ACCC) examining AI issues especially in pricing and online manipulation. There is discussion of AI-specific regulation and the Australian Human Rights Commission in 2021 recommended laws against AI discrimination, but concrete legislation is pending.

Other jurisdictions such as South Korea and India are in early stages as Korea has an AI Ethics Charter and is eyeing regulations, while India released an AI strategy but currently leans on IT and consumer laws to manage AI impacts.

The APAC region is thus a patchwork, since some countries lean toward EU-style comprehensive frameworks (e.g., Singapore's exploratory moves, China's targeted regulations), while others toward U.S.-style case-by-case enforcement.

5. A Common Trend: Diverging Tools, Converging Themes.

While regulatory strategies vary widely, several common principles are emerging across jurisdictions since i) transparency and disclosure of AI use in consumer interactions is becoming a baseline expectation,⁴⁵ ii) human oversight and recourse are widely endorsed,⁴⁶ iv) accountability and liability for AI outcomes are shifting toward the deployers and creators of AI,⁴⁷ iv) regulation is characterized by a risk-based tailoring.⁴⁸

⁴⁵ Whether by law or guidance, companies are increasingly pressed to inform consumers when an AI is driving contract terms or decisions, and to provide some intelligible explanation of those processes. This is evident in the EU (Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts OJ L 2024/1689 and Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation) OJ L119/1), Canada (AIDA proposal), China (algorithmic provisions), and in soft law like Singapore's framework.

⁴⁶ Nearly every jurisdiction emphasizes that AI should not be the sole and unchallengeable arbiter of significant consumer rights. Mechanisms for human intervention—be it internal officer review (as in Treasury Board of Canada Secretariat, 'Directive on Automated Decision-Making' (2019, amended 2022)), external appeal (as under GDPR/AI Act), or regulatory complaint routes—are being built or strengthened.

⁴⁷ Strict liability models (EU's forthcoming AI Liability regime, Projeto de Lei nº 2338/2023 (Brazilian Artificial Intelligence Bill)) and regulatory fine powers (EU AI Act, China's algorithm rules, potential FTC penalties) mean that businesses cannot hide behind the complexity of their tools. If an AI causes harm, the company using it will likely bear responsibility regardless of intent. This creates a strong incentive for diligent vetting and monitoring of AI behavior.

⁴⁸ Many jurisdictions differentiate AI applications by risk level. High-risk uses (especially those affecting life opportunities or fundamental rights) get more oversight, while low-risk uses are left largely alone. This risk lens is shared from the EU to Canada to Japan's principles, suggesting an emerging consensus that AI governance should not be one-size-fits-all but proportionate to impact.

However, divergences in implementation and scope persist. While the EU's horizontal, ex ante regulatory approach contrasts with the U.S.'s sectoral, ex post enforcement model, some countries lean heavily on ethical guidelines and industry self-regulation (Japan, Singapore), whereas others are embedding requirements in hard law (EU, China, Brazil). This fragmentation poses challenges for multinational businesses deploying AI in consumer contracts globally as they face inconsistent compliance obligations due to the fact that a practice acceptable in one country might be unlawful in another. It also opens the door to regulatory arbitrage, where companies might base certain AI operations in jurisdictions with laxer rules to avoid stricter oversight elsewhere (for instance, running AI analysis on consumer data in a country with no AI-specific laws, then using the results in stricter jurisdictions).

Efforts are underway to bridge these gaps, and international bodies are advocating for baseline standards such as the G7's Hiroshima AI Process that in 2023 called for interoperable principles on AI governance among major democracies, including common transparency and risk-management benchmarks,⁴⁹ or the OECD's AI Principles (adopted in 2019 and endorsed by 50+ countries) that similarly urge fairness, accountability, and transparency in AI and have informed laws like the EU's AI Act.⁵⁰ Furthermore, organizations like WIPO and UNCITRAL have begun exploring how AI intersects with international trade and contracting, including the potential for model contractual clauses or guidelines to ensure legal predictability across borders. While these initiatives are soft law, they indicate a converging recognition: despite divergent regulatory tools, the underlying policy goals for AI in consumer contracts are aligning globally around protecting consumers from opaque, biased, or unsafe AI practices without stifling innovation.

6. Liability, Remedies, and Enforcement.

A core concern in AI-powered consumer contracts is the delegation of legal decision-making to autonomous systems, and to mitigate this, emerging regulatory regimes mandate "human-in-the-loop" oversight for high-impact AI decisions (s.c. human oversight). Article 14 of the EU AI Act, for example, requires providers of high-risk AI systems to build in the ability for humans to monitor the AI's functioning and output, and to intervene or disable the system if it malfunctions or operates outside specified parameter. In the specific context of contracts, this might mean that whenever an AI proposes terms or makes a significant decision, like declining a loan or insurance coverage, a human official within the company must have the capacity to review and override that decision before it is finalized. The AI Act also calls for training those human overseers so they understand the AI's limitations and can spot when override is needed.

⁴⁹ G7 Leaders' Statement, *Hiroshima AI Process – International Guiding Principles for Organizations Developing Advanced AI Systems* (30 October 2023).

⁵⁰ OECD Council Recommendation on Artificial Intelligence, OECD/LEGAL/0449 (adopted 22 May 2019).

These requirements dovetail with existing laws on automated decisions, and GDPR Article 22 famously gives individuals the right not to be subject solely to automated decisions that have legal or similarly significant effects, unless they have given explicit consent or the decision is necessary for a contract and adequate safeguards (including human review) are in place. In practice, companies deploying AI contracting systems in the EU must offer consumers an option to request human reconsideration of an automated outcome. For instance, if an AI-generated rental agreement denies a person a lease, perhaps due to an algorithmic risk score of the applicant, the consumer should be informed of the AI's role and allowed to ask for a human to review the decision and the data underlying it. Similarly, Canada's proposed AIDA would require that consumers be told when they are subject to an automated decision and be given an avenue to seek human intervention as much as Brazil's draft AI law explicitly includes a right for consumers to request "clarifications from a human person" for any automated contractual term they deem abusive or incorrect, effectively ensuring human accountability.

The onus is increasingly on businesses to make these oversight and review mechanisms effective. Internally, companies are establishing AI review boards or designating specific staff who can interpret an AI's rationale, using tools like explainable AI outputs or model documentation, and veto the AI's decision if needed. Some firms are setting "confidence thresholds" for AI decisions, and if the AI's confidence in a particular output is below a certain level, it flags a human to make the call. Others have implemented periodic audits where a sample of AI-driven contracts are reviewed by humans to ensure the terms are fair and compliant. The aim of all these measures is to prevent AI from becoming a runaway train since no consumer should be bound by an automated decision that a reasonable human would find erroneous, unfair, or unlawful. By institutionalizing human oversight, regulators hope to catch the most egregious AI outputs before they cause harm and to give consumers reassurance that there is a "humane" fallback if they feel wronged by a machine decision.

On the other hand, effective consumer redress is critical to ensure accountability for AI-generated contracts (s.c. consumer redress mechanisms), and traditional contract remedies and consumer protection tools are being

adapted to the AI context, such as «contract avoidance or nullification,⁵¹ ii) compensation and damages,⁵² iii) regulatory orders,⁵³ iv) collective redress.⁵⁴

A recurring challenge, however, is accessibility of redress because if an AI system is making decisions at scale and speed, individual consumers might find it hard to pinpoint that something went wrong or to navigate the complaint process. Recognizing this, regulators have called for more user-friendly redress channels, and the European Data Protection Board, for example, has emphasized that rights like contesting an automated decision should be as accessible as the decision itself.⁵⁵ This has led to proposals for integrated redress mechanisms since AI-driven platforms could include built-in dispute resolution tools (chatbot assistants that help users file a complaint or appeal a decision, for instance). Some companies are experimenting with transparency portals where users can see a summary of how the AI scored or classified them and initiate a challenge if it seems incorrect, so that the advantage of these built-in systems is speed and scale as they can handle large volumes of minor disputes without requiring every consumer to go to court or a regulator.

Still, a tension exists between quick internal fixes and formal legal remedies, and companies might prefer to keep disputes in-house (to avoid precedent or

⁵¹ Consumers may have the right to cancel or rescind AI-generated contracts if consent was undermined by the AI's opacity or if the terms are fundamentally unfair. In the EU, a term generated by AI that was not transparently communicated could be deemed unenforceable under Council Directive 93/13/EEC of 5 April 1993 on unfair terms in consumer contracts OJ L95/29, art 6 (which provides that unfair terms are not binding on consumers) or treated as a misrepresentation allowing avoidance. National laws may also provide doctrines of error or vitiated consent – if a consumer can show they were mistaken about a contract's terms due to AI complexity, a court might void the contract.

⁵² Where an AI-driven contract causes quantifiable harm, consumers can seek compensation. This might be through general contract law (damages for breach or for entering into a contract under false pretenses) or through statutory remedies. For example, if an AI system systematically overcharged consumers beyond the agreed price, they could sue for restitution of the overcharges. Under data protection law like Regulation (EU) 2016/679 (GDPR) OJ L119/1, art 82, individuals can claim compensation for material or non-material damage resulting from unlawful processing—imagine an AI misuses personal data to set a price and the person suffers financially. In some jurisdictions, tort law (negligence or product liability) could also come into play if the AI is viewed as a product or service that was defective.

⁵³ Consumer protection agencies can act on behalf of consumers to correct AI-driven harms. In the EU, data protection authorities (DPAs) have the power to order companies to stop processing data in violation of *GDPR*—so if a contract AI is using personal data unlawfully, a DPA might order that system halted or its models deleted. Consumer law enforcers can similarly demand that companies cease using unfair algorithmic practices. For instance, a national consumer authority might issue an injunction against an online service to stop using an AI clause that automatically renews contracts without proper consent.

⁵⁴ As noted with Brazil, in many jurisdictions consumer redress can be collective. The EU is moving toward greater collective consumer action capability (Directive (EU) 2020/1828 of the European Parliament and of the Council of 25 November 2020 on representative actions for the protection of the collective interests of consumers and repealing Directive 2009/22/EC OJ L409/1 will allow qualified entities to bring collective actions for consumer law violations, which could include AI-related abuses). This means large-scale issues—say a million consumers all subjected to an unfair AI-drafted term—could be addressed in one proceeding with potentially significant remedies (contract nullifications, bulk refunds, etc.).

⁵⁵ European Data Protection Board, 'Guidelines on Automated individual decision-making and Profiling for the purposes of GDPR' (WP251 rev 01, endorsed 2018), noting data subjects' right to challenge automated decisions.

liability), whereas consumer advocates stress the need for external oversight to ensure fairness. The ideal emerging is a multi-tiered redress system that starts with easy internal challenge options, but always allows escalation to regulators or courts if the outcome remains unsatisfactory or if systemic issues are suspected.

Certainly, one of the most striking legal developments in response to AI in consumer contracts is the consideration of strict liability frameworks (s.c. strict liability regimes). Indeed, in traditional negligence-based liability, a consumer harmed by an AI decision would have to prove that the company deploying the AI failed to exercise due care, but given the complexity and opacity of AI, proving fault can be prohibitively difficult for an ordinary consumer. As a result, lawmakers are shifting the burden of risk to the side of the AI providers and users.

Under the EU AI Act, while primarily a regulatory statute, there is a hint of strict liability in its penalty scheme so that if a provider places a non-compliant high-risk AI system on the market (e.g., one that flouts the Act's requirements) and that results in harm, the provider faces administrative fines without need to show intent or negligence.⁵⁶ Moreover, the EU was concurrently working on an AI Liability Directive (proposed in 2022)⁵⁷ which will make it easier for victims of AI-related harm to seek compensation by, for instance, enabling courts to presume a causal link between an AI system's failure to meet legal requirements and the damage that occurred.⁵⁸ This is not pure strict liability in the sense of automatic liability for any harm, but it greatly alleviates the victim's burden of proof.

Brazil's Projeto de Lei nº 2338/2023 (Brazil's Proposed AI Regulation)⁵⁹ establishes a risk based framework for AI, focusing on ethics, transparency, and human right, explicitly leans toward strict liability since Articles 19–22 of the bill would hold AI system providers and operators jointly liable for damages caused by their systems, subject to limited defenses, such as proving the harm was caused by a force majeure or the victim's fault. The rationale is that those who create or profit from AI should internalize the costs of its failures, and

⁵⁶ See *Artificial Intelligence Act*, Art 71(3)(b).

⁵⁷ In September 2022, the European Commission introduced a Proposal for an Artificial Intelligence Liability Directive (AILD) with the objective of adapting the existing framework of non-contractual civil liability to harms arising from the use of AI systems. The proposal sought, inter alia, to mitigate the evidentiary difficulties typically faced by claimants in AI-related disputes by recalibrating traditional rules on the burden of proof. A central innovation in this respect was the introduction of a rebuttable presumption of causality, designed to facilitate the establishment of a causal link between the malfunction or unlawful operation of an AI system and the damage suffered. Notwithstanding these aims, the legislative initiative did not progress to adoption. In February 2025, the Commission formally withdrew the proposal, citing persistent disagreement among stakeholders as well as a broader policy shift towards streamlining and simplifying the EU's digital regulatory framework; subsequent institutional communications and reports have confirmed that the AILD proposal was definitively abandoned.

⁵⁸ European Commission, 'Proposal for a Directive of the European Parliament and of the Council on adapting non-contractual civil liability rules to artificial intelligence (AI Liability Directive)' COM (2022) 496 final, introducing burden-of-proof easements for victims of AI harm.

⁵⁹ The Bill has been approved by the Senate in December 2024, and now it is in the Chamber of Deputies.

therefore if an AI unintentionally generates a contract term that is illegal,⁶⁰ the lender using that AI could be liable to all affected borrowers for the illegal overcharge, even if the lender exercised some care in choosing the AI. The mere fact that the AI behaved in a forbidden way would trigger liability.

The benefits of strict or near-strict liability in this context are clear: it gives companies a strong incentive to ensure their AI systems are safe, fair, and compliant from the outset, since they will pay for any harm caused regardless of precautions, and it spares consumers the nearly impossible task of dissecting an AI's inner workings to prove a company's fault. It aligns with the idea of enterprise liability in products, and thus AI is akin to a complex product, and if it is defective, the manufacturer (or deployer) pays.

However, companies and some scholars have raised concerns that overly strict liability could chill innovation, especially for smaller enterprises, and, indeed, if any minor glitch could lead to open-ended liability, firms might avoid using AI in beneficial ways.⁶¹ Lawmakers are therefore calibrating these regimes by, for example, limiting the scope of strict liability to certain high-risk scenarios or capping damages in some cases. The AI Liability Directive proposal, for instance, did not impose strict liability per se, that might come in a separate Product Liability Directive update for AI, but it did reverse burdens of proof in crucial ways.

In sum, the trend is toward making AI providers the “insurers” of their technologies, and just as product liability law pushed manufacturers to improve product safety, knowing they would pay for injuries, AI liability rules aim to push developers to invest in bias mitigation, robustness, and thorough testing. For consumers, this offers greater protection, and, indeed, if they are wronged by an automated contract, they have a clearer path to remedy, without complex evidentiary battles over how an algorithm was coded or trained.

On another side, to ensure ongoing compliance and to detect issues proactively, regulators are increasingly requiring periodic audits of AI systems and transparency reports, and the EU AI Act, for instance, mandates that providers of high-risk AI keep detailed technical documentation and logs of the system's performance (accuracy, errors, bias incidents, etc.) and, where appropriate, submit to conformity assessments and post-market monitoring (Articles 61–62). Providers must report “serious incidents” and malfunctions to authorities within a set time frame, and in a consumer contracts context, a “serious incident” might be something like a systemic error in a contract AI that caused financial loss to many consumers or a detected pattern of illegal discrimination in the AI's outputs.

These audit and reporting obligations serve several purposes, and while internally, they force companies to maintain awareness of how their AI is

⁶⁰ For example, the generated contract inadvertently violates an interest rate cap leading to an usurious loan.

⁶¹ H Zech, ‘Liability for AI: public policy considerations’ (2021) 22 (1) ERA forum, 147 ff.

behaving in the field, not just in pre-launch testing,⁶² while externally, audits can reassure regulators and the public.⁶³

In the U.K., the Information Commissioner's Office (ICO) has recommended in its 2020 guidance "Explaining decisions made with AI" that organizations use techniques like impact assessments and independent verification to ensure their AI decisions are understandable and fair.⁶⁴ While not legally mandated for all sectors, such guidance sets expectations, and if a problem arises and it turns out a company never audited its AI, that lack of diligence could be used against it in enforcement or litigation.

Another aspect is transparency reporting, and early signs are already visible because the EU Digital Services Act (though not about contracts, but about platforms) requires big online platforms to explain their algorithms in reports. One could imagine future rules requiring a company that uses AI in consumer contracts to publish an annual transparency report outlining how many contracts were auto-generated, how many were challenged by consumers, what measures were taken to improve the AI's fairness, etc. Some large tech firms have voluntarily released AI transparency reports, preempting regulation.

Finally, regulatory sandboxes are being used to encourage companies to come forward and have their AI scrutinized in exchange for guidance rather than punishment,⁶⁵ so that a consumer-facing AI contract tool could be trialed in a sandbox to identify issues early. The data from these sandboxes (with appropriate confidentiality) can also inform regulators about what audit standards to set.

In summary, mandatory audits and reporting transform AI governance from a reactive to a proactive posture since, on the one side, they recognize that waiting for complaints or disasters (ex post enforcement) is insufficient, and, on the other, regulators want to know about potential problems before they scale. For businesses, this means firstly investing in compliance personnel who understand AI (much as data protection law led to the rise of Data Protection Officers), and secondly documenting everything, such as why a certain data set was used, how bias was tested, what thresholds are set for acceptable error rates, and so on. That documentation becomes both a shield (showing good faith efforts) and a sword (providing evidence if enforcement occurs).

Despite the strengthening of enforcement tools and remedies described above, there is growing recognition that after-the-fact remedies (s.c. ex post remedies) have inherent limitations in the context of fast-moving AI systems.

⁶² A common problem with machine learning systems is drift since over time, as input data changes, the AI's outputs might start violating assumptions. Continuous monitoring can catch these issues, for example, noticing that an AI's approval rates for certain groups are diverging.

⁶³ Some jurisdictions are considering requiring third-party audits for certain AI systems (similar to financial audits). An independent auditor could evaluate, for example, a lending AI's compliance with credit discrimination laws and produce a report.

⁶⁴ See, for further details, UK Information Commissioner's Office, 'Explaining decisions made with AI' (2020) <https://ico.org.uk/for-organisations/uk-gdpr-guidance-and-resources/artificial-intelligence/explaining-decisions-made-with-artificial-intelligence/> accessed 20 April 2026.

⁶⁵ For instance, the EU AI Act encourages Member States to set up sandboxes where companies can test innovative AI under regulator supervision. See Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts [2024] OJ L 1689/1, art 53.

In this regard, several challenges have been noted, moving from the speed and scale of harm. In fact, AI systems can affect thousands of consumers almost instantaneously, and by the time a single consumer realizes a term is unfair and pursues a complaint, the AI may have generated similar terms for many others. Enforcement actions and court cases typically take months or years, during which consumers continue to suffer harm.⁶⁶ The deterrent effect of ex post fines is also dampened if companies calculate that the profit from the AI practice outweighs the risk of a penalty down the line.

Then information and power asymmetry can be noted, and consumers often do not know they have been wronged by an AI, or they cannot prove it since AI decisions might be invisible.⁶⁷ This makes it hard to trigger ex post remedies, which usually rely on an individual complaint or dispute, and unlike a clear breach of contract, like non-delivery of goods, AI-driven inequities can be subtle and hidden, thus under-reported and under-remedied.

Further, procedural hurdles can emerge even when consumers do seek redress,⁶⁸ and regulatory agencies have finite resources and may prioritize only the most egregious cases. All this means many AI-related harms could slip through cracks.

Finally, as noted, proving causation and harm from an AI is complex (s.c. evidentiary complexity), and although strict liability helps, not all jurisdictions have it.⁶⁹

In light of these issues, there is a push towards ex ante safeguards (preventive measures) to reduce reliance on ex post fixes, and, apart the ones already discussed (i.e. mandatory audits, impact assessments, etc.), proposals include certification schemes or “quality seals” for AI systems (so consumers can trust certified systems), and real-time monitoring by regulators of high-impact AI (akin to how financial regulators monitor stock trading algorithms).

Furthermore, empowering consumers before harm occurs is key, and education and transparency can help here. If consumers know an AI is being used and have access to explanation tools, they might question decisions immediately, leading to quick corrections. Some companies have started offering “AI challenge” buttons – if you think a decision is wrong, click and it will be reviewed -, and this built-in feedback loop can catch errors early.

Another angle is public enforcement and collective mechanisms as regulators are considering more proactive auditing (not waiting for a

⁶⁶ For example, if an AI-driven subscription service sneakily shortens renewal notice periods (making it easy to trap consumers in renewals), an enforcement action might eventually impose fines and mandate clearer terms, but in the interim countless consumers might have paid fees they can't easily recover.

⁶⁷ For example, if a consumer is quoted a higher price than someone else for the same service, s/he likely will not know, as much as if an AI declines a consumer's application based on a biased logic, s/he may suspect unfairness but lack evidence.

⁶⁸ Arbitration clauses (often present in consumer contracts) might require disputes to be handled individually outside of court, limiting class actions. Companies may use NDAs in settlements to prevent disclosure of AI flaws.

⁶⁹ In a lawsuit, getting access to the AI's code or training data to demonstrate the fault might require expert testimony and discovery battles. Many individuals will not have the means to pursue that, effectively leaving them without remedy.

complaint).⁷⁰ In Europe, cooperation among data, consumer, and competition authorities is being fostered to tackle algorithmic issues from multiple angle, since an AI practice might simultaneously violate data law, consumer law, and antitrust, and, in general, joint task forces can lead to swifter, coordinated enforcement.

Finally, there is interest in technological solutions, and it is worth asking if “algorithmic insurance” can emerge:⁷¹ where companies insure consumers against certain AI errors? Or perhaps third-party services that scan your contracts and alert you to possibly unfair AI-generated terms (an AI to watchdog the AI, ironically).

In sum, while legal remedies after harm remain crucial, and are being strengthened, they are increasingly seen as the safety net, not the primary strategy. The focus is shifting toward preventing harm through design and oversight, recognizing that by the time a court case ends or a fine is levied, the AI landscape may have moved on and the consumers’ trust already undermined. The next section’s case studies illustrate both the use of ex post enforcement and its struggles in keeping pace with AI’s rapid deployment.

7. Case Law and Enforcement Trends in the European Union and in the U.S.A.: Enforcement Under the AI Act and GDPR, on the one side, FTC and CFPB Enforcement, on the other.

European regulators have begun to test their new powers on AI-driven contracting practices, even as the ink dries on the AI Act. A notable example arose in December 2025 when the European Commission has initiated a formal antitrust investigation to determine whether Google has infringed EU competition law by leveraging web publishers’ content, as well as content uploaded to the YouTube video-sharing platform, for artificial intelligence (“AI”) purposes. In particular, the Commission is examining whether Google has abused its dominant position by using web publishers’ content to power its generative AI search features (AI Overviews and AI Mode) without adequate remuneration or a genuine opt-out, and by training its generative AI models on YouTube content under mandatory, non-remunerated licensing conditions.⁷² The investigation also considers whether Google’s preferential access to such content, combined with restrictions preventing rival AI developers from using YouTube data, distorts competition to the detriment of publishers, creators, and competing AI providers, potentially infringing Article 102 TFEU and Article 54 EEA Agreement. What makes the case particularly significant is its potential to give rise to parallel enforcement under the GDPR. In addition to competition law concerns, data protection authorities – most notably those of Ireland and the Netherlands – may scrutinise whether Google’s use of publishers’ and

⁷⁰ For example, the FTC has publicly stated it will not hesitate to investigate companies on its own initiative if it suspects algorithmic bias or deception, even absent a specific complaint.

⁷¹ See, for further details, D Bertsimas and A Orfanoudaki, ‘Algorithmic insurance’ arXiv preprint arXiv:2106.00839 <https://arxiv.org/pdf/2106.00839> accessed 30 June 2026.

⁷² European Commission, ‘Commission opens investigation into possible anticompetitive conduct by Google in the use of online content for AI purposes’ (Press release, 9 December 2025) https://ec.europa.eu/commission/presscorner/detail/da/ip_25_2964 accessed 23 April 2026.

creators' content for generative AI purposes entails the processing of personal data in a manner incompatible with the GDPR. In particular, the large-scale ingestion and reuse of online and audiovisual content could involve profiling or other forms of personal data processing without a valid legal basis or effective consent, potentially engaging Articles 6 and 9 GDPR alongside the antitrust assessment.

This layered enforcement logic illustrates the complementary operation of EU regulatory instruments: the AI Act governs the design, deployment, and governance of algorithmic decision-making systems, while the GDPR regulates the lawfulness of the personal data processing underlying those systems.⁷³ For undertakings, this creates cumulative exposure, as the same AI-enabled practice may trigger parallel proceedings and sanctions by different authorities. More broadly, it exemplifies the systemic interoperability of EU law, whereby a single conduct may simultaneously raise issues under AI regulation, data protection law, competition law, and potentially consumer protection or unfair contract terms rules. Although such multi-track enforcement is resource-intensive, it enables regulators to address the full spectrum of harms generated by complex AI-driven practices.

Another emerging pattern in EU enforcement is cooperation through the European Data Protection Board (EDPB) and the new European Artificial Intelligence Board, a key advisory body that was created by the AI Act, since these bodies enable joint investigations and consistent application of rules. The streaming service case saw cross-border cooperation, which is something pioneered in GDPR enforcement via the one-stop-shop mechanism and likely to be mirrored in AI Act enforcement.

It is worth noting that as of 2026, most AI Act enforcement actions are at an early stage, since the Act's provisions on high-risk systems will fully apply only after a transitional period, so precedent is limited. However, regulators have signaled they are not waiting idly, and many Data Protection Authorities in the EU have already issued opinions that automated consumer contract decisions should be considered "high risk processing" under GDPR requiring Data Protection Impact Assessments.⁷⁴ Failing to do such an assessment, which overlaps with AI Act requirements, can itself lead to fines. In this sense, a comparable enforcement outcome is the 310 million euros fine against LinkedIn by the Irish Data Protection Commission in 2024 for violations of the GDPR relating to behavioural analysis of users' data, underscoring how EU data protection authorities are prepared to sanction exploitative data processing practices tied to algorithmic profiling and targeting.⁷⁵

⁷³ See, for further details, A Fiorentini, 'Rethinking the Institutional Framework for the EU Data Regulations: Towards an Enhanced Administrative Integration' (2025) 18(3–4) *Rev Eur Adm Law* 51 ff.

⁷⁴ See, on this regard, A Galandarli, 'Mitigating AI risks: A comparative analysis of Data Protection Impact Assessments under GDPR and KVKK' (2025) 7(3) *J Data Prot & Privacy* 252 ff.

⁷⁵ See Irish Data Protection Commission, 'Irish Data Protection Commission fines LinkedIn Ireland €310 million' (24 October 2024) <https://www.dataprotection.ie/en/news-media/press-releases/irish-data-protection-commission-fines-linkedin-ireland-eu310-million> accessed 17 April 2026.

On the other side, U.S. enforcement authorities, though lacking AI-specific statutes, have been creative in applying existing laws to AI-related contract issues.

The FTC in particular has taken the lead in policing deceptive or unfair use of AI in consumer interactions. In the case *Federal Trade Commission vs. MediaAlpha*, mentioned earlier, the Assurance IQ, LLC and MediaAlpha, Inc. will pay a total of \$145 million to settle Federal Trade Commission charges that they misled millions of consumers seeking to buy comprehensive health insurance. In two separate actions, the FTC alleged that both Assurance and MediaAlpha deceived consumers and led them to purchase plans that did not provide the promised health care coverage, and bombarded consumers with telemarketing and robocalls.⁷⁶ The Federal Trade Commission initiated an enforcement action against MediaAlpha, Inc. and its wholly owned subsidiary, QuoteLab, LLC, alleging unfair and deceptive acts or practices in violation of § 5(a) of the Federal Trade Commission Act, 15 U.S.C. § 45(a), as well as contraventions of the Telemarketing Sales Rule, 16 C.F.R. Part 310, and the Government and Business Impersonation Rule, 16 C.F.R. Part 461. The complaint concerned the defendants' operation of a digital lead-generation ecosystem in the health insurance market that systematically misrepresented both the nature of the offered products and the source of the solicitation.

According to the Commission, the defendants disseminated online advertisements and operated web interfaces that conveyed false indicia of governmental endorsement or affiliation – particularly with federal health insurance programs – while advertising ostensibly comprehensive, low-cost health insurance coverage. In practice, these representations functioned as inducements to extract sensitive consumer data, which were subsequently monetised through the sale of leads to third-party telemarketers marketing private insurance products that materially diverged from the advertised coverage.

The proceeding was resolved through a stipulated order imposing a permanent injunction and a monetary judgment of 45 million U.S. dollars. The settlement underscores the FTC's increasingly assertive regulatory posture toward deceptive digital marketing and lead-generation practices in the insurance sector, particularly where such practices exploit information asymmetries and impersonation of public authorities to distort consumer decision-making.

The Consumer Financial Protection Bureau (CFPB), meanwhile, has focused on ensuring AI does not become a loophole to evade fair lending and disclosure laws. In late 2023, the CFPB launched investigations into multiple fintech lenders employing AI algorithms for credit decisions. Although no public enforcement action has been concluded as of this writing, CFPB officials have issued guidance reinforcing that adverse action notices (required by ECOA and

⁷⁶ US Federal Trade Commission, 'Assurance IQ and MediaAlpha to Pay a Total of \$145 Million to Settle FTC Charges That They Misled Consumers Seeking Health Insurance' (7 August 2025) <https://www.ftc.gov/news-events/news/press-releases/2025/08/assurance-iq-mediaalpha-pay-total-145-million-settle-ftc-charges-they-misled-consumers-seeking> accessed 23 April 2026.

Regulation B when credit is denied) must be specific even if an AI is used.⁷⁷ One lender received a stern warning after a supervision exam found that their AI model occasionally denied credit based on complex patterns the lender could not explain to consumers as it basically issued notices saying that the application was denied due to proprietary algorithm. The CFPB made clear this is unacceptable and creditors cannot use the excuse of a “black box” to flout the obligation to tell consumers the main reasons for denial. That stance was formalized in a CFPB Circular in 2022 and reiterated in 2023 where it was reported established that lenders shall provide accurate reasons for adverse decisions, even if it means investing in interpretable AI or post-hoc explanation tools.⁷⁸

Additionally, the U.S. Department of Justice has shown interest in cases where AI may lead to discriminatory outcomes in violation of laws like the Fair Housing Act or Title VII (employment discrimination).⁷⁹ While those are not strictly “consumer contracts,” they intersect with automated decision-making in housing rentals, mortgage approvals, etc., and, in fact, a notable case in 2022 involved a landlord whose tenant screening algorithm deny rental housing based because it exhibit systemic inaccuracies, including: i) adverse information erroneously attributed to individuals other than the applicant, ii) outdated or legally obsolete data, and iii) inaccurate or misleading entries concerning arrests, criminal records, and prior evictions, which are often neither corrected nor expunged from the screening databases, inducing DOJ to intervene to enforce fair housing laws.⁸⁰ This indicates that when AI affects civil rights, U.S. authorities will adapt existing legal tools to address it.

In summary, U.S. enforcement to date is characterized by retrofitting old laws to new tech, whose advantage is flexibility as agencies can tackle egregious AI abuses without waiting for new legislation. The drawback is, on the one side, uncertainty because companies shall infer from these actions what is expected, and, on the other, reactive since regulators need to find out about the problem, often from complaints or reports, rather than having companies obliged to self-report as in the EU. Still, given the high profile of AI issues, the FTC and CFPB have been increasingly proactive, monitoring industry developments, issuing warnings, such as the FTC’s 2020 report “AI and Algorithmic Bias” and 2021 blog post “Aiming for truth, fairness, and equity in your company’s use of AI”, and collaborating with other agencies (e.g., the joint

⁷⁷ See Consumer Financial Protection Bureau, ‘Consumer Financial Protection Circular 2023-03: Adverse Action Notification Requirements in Credit Decisions Based on Complex Algorithms’ (19 September 2023) https://files.consumerfinance.gov/f/documents/cfpb_adverse_action_notice_circular_2023-09.pdf accessed 30 June 2026.

⁷⁸ See Consumer Financial Protection Bureau (n 17).

⁷⁹ See, for further details, N Babazadeh, A Washington and T Brown, ‘Civil Rights in the Digital Age: The Intersection of Artificial Intelligence, Employment Decisions, and Protecting Civil Rights’ (2022) 70(1) J Fed Law & Prac 75 ff <https://www.justice.gov/file/1189116/dl?inline=> accessed 30 June 2026.

⁸⁰ L Karpinski, ‘The Discriminatory Impacts of AI-Powered Tenant Screening Programs’ Georgetown J Poverty Law & Pol <https://www.law.georgetown.edu/poverty-journal/blog/the-discriminatory-impacts-of-ai-powered-tenant-screening-programs/> accessed 30 June 2026.

FTC/CFPB/DOJ/EEOC statement on AI in April 2023 emphasizing a united front against discrimination).⁸¹

8. United Kingdom and The Algorithmic Insurance Discrimination.

The United Kingdom, post-Brexit, is charting its own path on AI regulation, but its regulators have remained vigilant about algorithmic harms under existing laws. A high-profile instance in 2023 involved the Financial Conduct Authority (FCA) examining the use of AI in pricing auto insurance. The FCA had been alerted by consumer groups, like Citizens Advice, to a possible “ethnicity penalty” in car insurance pricing since FCA’s analysis finds links between the price of motor insurance and local area ethnicity in England and Wales.⁸² The FCA’s investigation found that some insurers used AI-driven pricing models that incorporated geographic data (postcodes) and proxy variables that, unintentionally or not, resulted in minority-dense neighbourhoods being quoted higher premiums for the same risk profile. While insurers have long used location as a factor (due to crime rates, etc.), the AI models were complex and included correlations that went beyond traditional actuarial factors.⁸³

The FCA did not accuse insurers of intentional racial discrimination, but it raised concerns that the models could breach the Equality Act 2010 if they could not be justified as proportionate to risk. In response, the FCA issued new guidance to the insurance industry, effectively amounting to enforcement by guidance, and it instructed insurers to audit their algorithms for potential bias and to remove any factors that cannot be transparently defended. It also highlighted that under the FCA’s Treating Customers Fairly (TCF) principle,⁸⁴ outcomes that systematically disadvantage a group of consumers could be deemed unfair. Several insurers, facing this pressure, voluntarily revamped their models, and one large insurer suspended the use of an external AI pricing tool entirely until it could be assured the tool’s decisions were explainable and free of proxy bias.⁸⁵

This case illustrates a more collaborative style of enforcement, considering that the FCA did not need to levy fines, but the reputational and compliance pressure was enough to prompt change. It also underscores the importance of

⁸¹ T McSweeney, LJ Tiedrich and J Pierce, ‘FTC Provides Guidance on Use of AI and Algorithms’ (2020) 3(6) J Robotics Artif Intell & L 431 ff; E Jillson, ‘Aiming for truth, fairness, and equity in your company’s use of AI’ (Federal Trade Commission, 2021) <https://privacysecurityacademy.com/wp-content/uploads/2021/04/Aiming-for-truth-fairness-and-equity-in-your-companys-use-of-AI.pdf> accessed 30 June 2026.

⁸² G Reynolds, ‘Motor insurance pricing and local area ethnicity’ (Financial Conduct Authority, 2025) <https://www.fca.org.uk/news/blogs/motor-insurance-pricing-local-area-ethnicity> accessed 30 June 2026.

⁸³ See M Van Bekkum, F Zuiderveen Borgesius and T Heskes, ‘AI, insurance, discrimination and unfair differentiation: an overview and research agenda’ (2025) 17(1) Law Innovation & Tech 177 ff.

⁸⁴ A Georgosouli, ‘The FSA’s Treating Customers Fairly (TCF) initiative: What is so good about it and why it may not work’ (2011) 38(3) J Law & Soc 405 ff; D Millard and CJ Maholo, ‘Treating customers fairly: A new name for existing principles?’ (2016) 79 THRHR 594 ff.

⁸⁵ DIGWATCH, ‘UK watchdogs warned over AI risks in financial services’ (21 January 2026) <https://dig.watch/updates/uks-warned-ai-risks-in-financial-services> accessed 23 April 2026.

explainability as the insurers themselves reportedly struggled to fully explain how the AI set prices, making it hard to contest allegations of bias. The lesson learned was that insurers, and by extension, other businesses using AI in contracting, shall be able to open the hood of their AI models and demonstrate that they operate within legal bounds, otherwise, they risk regulatory intervention or loss of consumer trust.⁸⁶

The U.K. is also in the process of developing its AI regulatory framework. In 2023, the U.K. government published a white paper proposing a principles-based approach (with principles like safety, transparency, fairness) to be implemented by sectoral regulators rather than via a single AI law. The FCA's actions on insurance AI can be seen as a testbed for that approach, using existing powers to enforce emerging principles, and it can be expected UK regulators in other sectors (e.g., the Competition and Markets Authority or the Information Commissioner's Office) to take on similar roles, scrutinizing AI practices within their remit.⁸⁷

Therefore as of early 2026, the U.K. has not enacted a single, comprehensive legislative framework for artificial intelligence analogous to the European Union's AI Act. Instead, the U.K. has adopted a sectoral, "pro-innovation" regulatory approach, in which existing authorities, such as the Information Commissioner's Office (ICO), the Competition and Markets Authority (CMA), and the Financial Conduct Authority (FCA) —apply five guiding principles to AI systems: safety, transparency, fairness, accountability, and contestability. While the Artificial Intelligence (Regulation) Bill 2024-25 seeks to establish statutory oversight and create a dedicated AI Authority, the U.K. government has predominantly favored a flexible, principles-based framework over centralized, prescriptive legislation, in contrast to the EU's risk-based AI Act, which is scheduled to enter into full force by August 2026.

9. A Hypothesis of Multi-Jurisdiction Litigation: Dynamic Pricing Clause and Emerging Enforcement Patterns.

The borderless nature of AI-driven commerce can lead to fascinating multi-jurisdictional legal conflicts. A hypothetical but instructive scenario could involve, for example, a global streaming service (let us call it StreamCo) that implemented an AI-driven dynamic pricing clause in its subscription contracts. The clause would allow StreamCo to adjust the monthly subscription fee for individual users every six months based on an AI assessment of the user's engagement, willingness to pay, and alternative options. Essentially, loyal customers who rarely switched services might see price hikes, whereas price-

⁸⁶ S Martinelli, 'AI as a tool to manage contracts consequences and challenges in applying legal tech to contracts management' (2023) 31(2-3) ERPL 411 ff; M Herbosch, 'To err is human: Managing the risks of contracting AI systems' (2025) 56 Comput Law & Secur Rev 106110.

⁸⁷ For example, the CMA in 2022 began analyzing online choice architecture and algorithms for potential consumer manipulation, which could lead to actions if companies are found to use AI-driven dark patterns in contracts or sales. See Competition & Markets Authority, *Online Choice Architecture: How digital design can harm competition and consumers* (London 2022) https://assets.publishing.service.gov.uk/media/624c27c68fa8f527710aaf58/Online_choice_architecture_discussion_paper.pdf accessed 30 June 2026.

sensitive customers might get discounts to prevent churn. This practice would be deployed worldwide via the same AI system, and this dynamic pricing clause would face legal challenges in multiple countries almost simultaneously as firstly in the EU, consumer organizations would file complaints alleging that the clause violated the transparency and fairness requirements of EU consumer law. The claim would be that users were not adequately informed about how their price was being set (lack of transparency) and that it constituted a potentially unfair contract term under Directive 93/13/EEC, because the company had unilateral power to change price based on vaguely defined criteria. Data protection issues would also be raised, since the AI used personal data (viewing habits, device information, etc.) to make pricing decisions without clear consent for that purpose (implicating GDPR's purpose limitation). An EU consumer protection authority ultimately could find the clause unenforceable and could order StreamCo to cease using it and to revert impacted customers to a stable pricing model, citing that genuine informed consent to such a dynamic scheme was not obtained.

In the United States, a class-action lawsuit would be filed in California under state consumer protection and contract laws. However, U.S. law (absent specific regulation) tends to give businesses more leeway for variable pricing as long as it's disclosed in the contract terms. StreamCo would argue that its Terms of Service explicitly allowed price changes and that continued use of the service after notice constituted acceptance. The U.S. court could be sympathetic to the argument that no law forbids individualized pricing and that consumers had the freedom to cancel if they did not like the new price. The case would be further complicated by arbitration clauses, according to which StreamCo moved to compel arbitration per its contract. Eventually, the class-action would not be certified and some claims would be to individual arbitrations, with mixed outcomes. There would not be a broad U.S. regulatory action, as neither the FTC nor state AGs would take issue with the practice provided no protected class discrimination was involved.

In Singapore, a group of consumers would bring the issue to the consumer tribunal (a sort of small-claims court for unfair practices) according to Singapore's consumer law (Consumer Protection – Fair Trading – Act) that allows for declarations of unfair practices. The tribunal would find the communication of the price changes to be lacking clarity and would order StreamCo to provide refunds to a set of consumers who would have complained, on the basis that they were not clearly notified of the AI-driven adjustment. However, the tribunal would not ban the practice outright, but it suggested that StreamCo could continue dynamic pricing if it improved transparency, for example, using an upfront notice like «your price may increase or decrease based on your usage patterns» in bold text during sign-up.

This scenario exemplifies regulatory fragmentation since the same business model can be simultaneously legal, illegal, or conditionally acceptable depending on the jurisdiction. For companies, this raises the question: do they tailor their AI systems country-by-country (geofencing certain features), or do they adhere to the strictest common denominator globally? StreamCo in this scenario decided to apply the EU's standard globally (i.e., they would drop the

dynamic pricing worldwide) because the complexity of maintaining different models was high and the EU outcome would threaten large fines if ignored.

It also highlights the risk of forum shopping by consumers and regulators as where laws are stricter, action will be taken that can influence global operations. And conversely, companies might choose to roll out contentious features first in places where they expect lighter touch oversight.

Finally, the need for mutual recognition or cooperation is underscored because if regulators could coordinate, perhaps a common stance on dynamic pricing could emerge to avoid such divergent results. Some effort in that direction is seen in international forums (as mentioned, G7, OECD), but it remains early. Until then, multinational companies employing AI in contracting face a puzzle of complying with a jigsaw of rules, often defaulting to the toughest ones as a safe harbor.⁸⁸

Across these examples, several enforcement patterns in AI-powered contract oversight are emerging,⁸⁹ such as disclosure and consent failures, which are low-hanging fruit as the easiest target for regulators is when companies do not even meet basic transparency. Many early enforcement actions (EU streaming service fine, FTC MediaAlpha case, etc.) revolve around companies not telling consumers what the AI is doing or misrepresenting it. Regulators consistently penalize hiding the ball, whether it is failing to disclose an AI's involvement or failing to explain outcomes. This suggests companies should prioritize clear communication as a first line of defense.

Secondly, as seen, enforcement often happens under multiple legal frameworks at once (multiple laws in play) and a single AI practice might trigger data protection law, consumer contract law, anti-discrimination law, and more. Agencies are increasingly collaborating (e.g., joint statements, coordinated raids), and for businesses, it means that a holistic compliance approach is needed, fixing an issue for GDPR but not considering consumer law could still land a company in trouble. As a consequence, one can expect to see more "joint enforcement actions" and cross-referrals between authorities, like a consumer authority referring algorithm issues to the privacy commissioner, and vice versa.

Thirdly, enforcement outcomes are not just punitive (fines) but often corrective as remedies include forward-looking measures. Settlements and orders frequently require companies to implement compliance programs, conduct audits, or even algorithmic changes.⁹⁰ These forward-looking remedies

⁸⁸ HW Liu and CF Lin, 'Artificial intelligence and global trade governance: a pluralist agenda' (2020) 61 Harv Int'l LJ 407 ff; IHY Chiu and EW Lim, 'Managing Corporations' Risk in Adopting Artificial Intelligence: A Corporate Responsibility Paradigm' (2021) 20 Wash U Global Stud L Rev 347 ff.

⁸⁹ S De Spiegeleire, M Maas and T Sweijs, *Artificial Intelligence and the Future of Defense: Strategic Implications for Small-and Medium-Sized Force Providers* (The Hague 2017) 61 ff <https://hcass.nl/wp-content/uploads/2017/05/Artificial-Intelligence-and-the-Future-of-Defense.pdf> accessed 30 June 2026; A Reddy and others, 'Survivable AI for Defense Strategies in Industry 4.0' in M Thirunavukkarasan and others (eds), *Artificial Intelligence and Machine Learning for Industry 4.0* (Wiley 2025) 169 ff; J Schuett, 'Three lines of defense against risks from AI' (2025) 40(2) AI & Society 493 ff.

⁹⁰ For instance, FTC settlements might mandate an independent algorithm audit annually for 20 years (as they have done in privacy cases).

indicate regulators want to fix the root cause, not just punish past behavior, and in the U.K. insurance case, effectively the remedy was “improve your model or stop using it.” So enforcement is steering the development of AI tools toward better practices.

Fourthly, there is a shift from waiting for consumer complaints to regulators using their own investigative powers or market studies to find problems (s.c. proactive detection).⁹¹ Since the FCA did research on insurance pricing, the CMA is studying algorithms, the CFPB is analyzing lending data for patterns, this proactive approach is necessary because consumers often cannot raise the alarm on invisible algorithmic issues. It means companies might be scrutinized even if no one has yet complained and agencies might knock on the door asking to audit your AI processes.

Fifthly, a strong stance in one jurisdiction can influence behavior elsewhere, even without formal cooperation (s.c. global ripple effects).⁹² After the EU’s actions on the streaming service, other companies preemptively reviewed their AI contract terms globally as much as after the FTC’s AI warnings, some firms chose to implement explainability features not just in the U.S. but worldwide, to show leadership and avoid being the next target. This unofficial harmonization by deterrence is interesting: savvy companies realize that being a responsible AI actor could become a competitive advantage (or at least save them legal headaches), so they voluntarily adhere to higher standards than what their local law strictly requires.

In conclusion of the case law survey, it is evident that AI in consumer contracts is no longer a hypothetical concern for enforcers, but it is here and now. While the legal theories used might be old (unfair practices, lack of consent, discrimination), they are being applied to new fact patterns with vigor and companies must stay abreast of this evolving jurisprudence, as it effectively creates the contours of AI contract law even before all the statutes are in place.

10. Toward an emerging Global Convergence?

Despite the absence of a unified global framework and the profound structural differences, a set of converging legal principles seems to be emerging across jurisdictions for regulating AI in consumer contracts, such as transparency and explainability.⁹³ There is a broad consensus that consumers should be informed when AI is involved in contracting and that significant automated decisions should be explainable. The EU AI Act hardcodes this and high-risk AI shall have accompanying information and technical explainability (Article 6), Canada’s AIDA proposal includes transparency requirements, and even U.S. regulators, via guidance, emphasize not hiding AI usage. Many

⁹¹ A Mahmutovic, ‘The EU AI Act: a proactive framework for comprehensive AI regulation’ (2025) 33 Int J Law & IT eaf028.

⁹² E Haber, ‘The Invisible Ripple Effect’ https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5837642 accessed 30 June 2026.

⁹³ N Balasubramaniam and others, ‘Transparency and explainability of AI systems: From ethical guidelines to requirements’ (2023) 159 Information and Software Technology 107197.

jurisdictions now expect that if an AI sets a price or term, consumers have a right to know that and, at least upon request, to get an explanation of the main factors. This principle is reflected in the OECD AI Principles as well, advocating transparency and responsible disclosure.

Building on data protection regimes, a nearly universal norm is the right for individuals to contest automated decisions that seriously affect them (consumer challenge and consent rights). The exact mechanics vary (GDPR's Article 22 vs. a more general consumer complaint avenue in other places), but the idea is the same as AI should not have the final say unilaterally. Moreover, obtaining meaningful consent for AI-driven practices is emphasized.⁹⁴ For example, dynamic personalization of contract terms arguably requires consent (or at least clear notice and opt-out) in many frameworks, and it may be seen this in proposed laws and in regulators' expectations (e.g., the notion of "no AI exemption" to consent requirements under privacy laws).

As discussed, regulators worldwide champion the idea that appropriate human oversight is essential (human accountability and oversight). This ranges from mandatory human review of certain AI outputs to holding a specific person accountable (Singapore's Model AI Framework suggests having an AI accountability function, while some proposals call for licensed AI professionals akin to how financial firms have compliance officers). The underlying principle is that AI should augment, not replace, human responsibility, and if an AI contract system causes a legal violation, there should be an identifiable human or corporate actor to answer for it.⁹⁵

Global discourse agrees that regulatory intensity should correlate with the risk posed by the AI application (risk-proportionate governance).⁹⁶ High-risk uses, like deciding who gets a mortgage, warrant close regulation, while low-risk uses, like auto-completing a mailing address in a form contract, can be lightly governed. This principle is explicit in the EU AI Act and implicitly followed elsewhere, and it is also reflected in soft law since the G20 and others have endorsed a risk-based approach as a rational way to regulate AI without stifling innovation in benign applications.

Finally, another principle taking hold is that companies should document their AI systems and be ready to demonstrate compliance, akin to "show your work" (accountability through documentation and audit).⁹⁷ This is seen in EU's requirement for technical documentation, in Canada's and Singapore's encouragement of keeping records of AI decision processes, and even in U.S. proposals for algorithmic accountability reports. The principle ties into

⁹⁴ J De Matas, J Wang and V Gupta, 'Narrative Transparency in AI-Driven Consent' (2025) 25(4) *Am J Bioethics* 136 ff.

⁹⁵ M Herbosch, 'To err is human: Managing the risks of contracting AI systems' (2025) 56 *Comput Law & Secur Rev* 106110.

⁹⁶ R Gangavarapu, *Mastering AI Governance: A Guide to Building Trustworthy and Transparent AI Systems* (Cham 2025) 121 ff; T Szadeczky and Z Bederna, 'Risk, regulation, and governance: evaluating artificial intelligence across diverse application scenarios' (2025) 38(1) *Security J* 35 ff.

⁹⁷ F Königstorfer and S Thalmann, 'AI Documentation: A path to accountability' (2022) 11 *J Responsible Tech* 100043.

corporate governance as boards and senior management should be accountable for how AI is used in their organization.

Crucially, these principles are being reinforced through international cooperation, and, for instance, the Global Privacy Assembly (an international network of privacy regulators) has issued declarations that echo these points, and the new Global Partnership on AI (GPAI) provides a forum for policy experimentation aligned with them.

Notwithstanding the scenario just described, it is worth noticing that divergence in implementation and scope still exists, since while agreeing on principles is one thing, implementing them is another. Different jurisdictions still exhibit significant divergence in legal approaches and scope of AI regulation, and this fragmentation means compliance for global companies is complex. It also means consumers have uneven protection depending on where they live. A European consumer might have multiple avenues to challenge an AI decision (under AI Act, GDPR, consumer law), whereas an American consumer might have to rely on general contract law and hope a court sees an AI-generated unfair term as unconscionable, which is possible but not guaranteed. A Brazilian consumer could benefit from strict liability and collective actions, while a consumer in, say, Indonesia currently might find no AI-specific recourse at all. Furthermore, significant differences remain in the underlying contractual traditions of the jurisdictions considered. In particular, divergences can be observed in the interpretation and role of principles such as good faith, substantive fairness, and contractual autonomy, as well as in choice of law rules and available remedies. These persisting differences suggest that, while some convergence may be emerging at the level of regulatory objectives, the implementation of a uniform global framework for AI consumer contracts remains premature.

Regulatory arbitrage is a real risk in this scenario, and companies might deploy more aggressive AI features in markets with looser rules⁹⁸ as much as data might be processed in jurisdictions with weaker data protection or AI oversight and then used to influence contracts in stricter jurisdictions and cloud services make such geographical shifting easy. If one country's laws ban a certain AI practice, but another's do not, a global company might simply geofence that feature, which raises questions about fairness: why should consumers in one country get better protections than another?

Reasonably, businesses will naturally seek to minimize compliance costs and liability, and the global AI regulatory patchwork offers opportunities to do so. Therefore, as an example, a company could base its AI R&D and model training in a country with lenient data rules (to freely scrape and use personal data), then export the trained model for deployment globally. If caught, the company might argue the harmful processing happened outside the jurisdiction of a stricter law. This raises challenges for regulators: can they assert jurisdiction extraterritorially? The GDPR, for instance, can reach data controllers abroad if they target EU individuals. The AI Act will similarly have extraterritorial reach for providers placing systems in the EU market, but enforcement across borders appears tricky. Secondly, firms might experiment with high-risk AI

⁹⁸ For example, a company could decide to implement dynamic pricing algorithms in North America but not in Europe, to avoid EU compliance burdens.

applications in countries without specific AI regulations to gather data and refine the system, and only later introduce them in regulated markets claiming the system has a proven safety record (even if that “testing” potentially harmed consumers elsewhere who had fewer protections). Then, there is also the phenomenon of forum shopping in litigation. If a global AI causes harm, plaintiffs will seek the venue with the most favorable law or precedent. Conversely, companies might include forum selection clauses and arbitration agreements to funnel disputes to jurisdictions or processes where the law of AI is underdeveloped or their liability is limited. Furthermore, one concrete area of concern is data flows: AI systems in contracts often rely on personal data. The EU’s stringent rules on data exports (Schrems II ruling, etc.) mean some companies might keep EU data in Europe, but use more data in the U.S. or Asia for building models. This could lead to inequities in how advanced or personalized the services are in different regions, essentially an arbitrage of data freedom. Additionally, with AI models like LLMs, once a model is trained, it can be deployed anywhere. If one country bans a certain type of AI (say deepfake-based contracting or AI that exploits behavioral biases), the model might still be accessible via the internet. A clear example is linked to applications: one country’s ban often just means users route around it via VPNs. So unilateral regulations can be undermined by the global availability of AI tools, unless there’s coordination.

On the enforcement side, lack of harmonization can also lead to inconsistent outcomes as it was shown in the StreamCo scenario, so that companies may face penalties in one place and not in another for similar conduct, which can be perceived as arbitrary and undermine the sense of justice or effective deterrence.

However, an interesting development in bridging legal gaps is the idea of mutual recognition of AI oversight and enforcement among jurisdictions (mutual recognition pilots), and consequently if one regulator takes an action or certifies a system, others could accept that outcome rather than re-inventing the wheel. It may be seen as some seeds of this, as the EU has been discussing cooperation agreements with like-minded countries (e.g., Canada, Japan) to recognize each other’s conformity assessments for AI systems. In practical terms, if an AI system is certified as compliant with the EU AI Act by an EU notified body, Canada might agree to accept that certification as evidence of meeting Canadian requirements (and vice versa). This is analogous to how some jurisdictions recognize each other’s product safety testing to ease trade. A concrete pilot is the EU-Canada Digital Cooperation initiative, as Canada’s Directive on Automated Decision-Making (for government) is one of the most advanced of its kind. Suppose a Canadian government service’s AI passes the Directive’s Algorithmic Impact Assessment at the highest level, and if that same service were offered to EU citizens, the EU could treat that AIA as partially satisfying EU AI Act requirements. In 2024, as mentioned, the EU and Canada were reportedly exploring an arrangement whereby a consumer can file a complaint about an AI system through either Canadian or EU channels and have it addressed collaboratively. While not yet formalized, the idea would be a step toward trans-Atlantic redress: a European using a Canadian AI service could complain to their local EU consumer authority, which would liaise with Canadian

counterparts to resolve it, and the remedy (say compensation or correction) would be recognized and enforced in both places. On enforcement, there are existing frameworks among some regulators for cooperation, like the Global Privacy Enforcement Network (GPEN) for data protection and the International Consumer Protection Enforcement Network (ICPEN) for consumer law. They could be extended to AI issues like a coordinated international sweep of e-commerce sites checking for AI-generated dark patterns, with enforcers in multiple countries acting in concert.

However, mutual recognition faces challenges, since countries need to trust each other's regulatory rigor, and there is politics involved. But it has precedence in financial regulation and data protection (the concept of "adequacy" in data transfers, for instance, is a kind of recognition of another's regime).⁹⁹ If, say, the EU deems a country's AI rules "equivalent," it might simplify compliance for companies operating in both.

For consumers, mutual recognition could mean no matter where an AI provider is based, they have a local avenue for recourse and that their issue will be handled with the same seriousness as if it occurred domestically.

Given the trends above, what strategies could drive the world toward a more harmonized governance of AI in consumer contracts? Scholars and policymakers have proposed several, but overall, the drive toward convergence does not mean every jurisdiction will be identical. In this perspective, several solutions have been proposed. First, Baseline International Standards: develop something akin to the Vienna Convention on Contracts for the International Sale of Goods (CISG), but for AI-related aspects of contracts. This could cover definitions (what is an "AI agent" legally? when is an automated contract valid?), obligations (like a duty to disclose AI use), and perhaps liability principles. While a full treaty may be far off, even a non-binding UN General Assembly resolution outlining baseline principles could guide national laws.¹⁰⁰ Second, Model Laws and Clauses: as UNCITRAL often does, create model legal provisions that countries can adopt or companies can insert into contracts. For example, a model clause might state: «Where one party uses an automated system to negotiate or form this contract, that party bears responsibility for any errors or omissions of the system, and shall inform the other party of the use of such system». If many contracts include such language, it could create a de facto global standard that courts enforce, even absent legislation.¹⁰¹ Third, Cross-border Redress Mechanisms: perhaps building on existing frameworks like the e-Courts or international arbitration, create specialized arbitration or ombudsperson schemes for AI disputes. One idea is an international ombuds

⁹⁹ See, for further details, M Andenas and IH-Y Chiu, *The Foundations and Future of Financial Regulation: Governance for Responsibility* (Routledge 2014), where the A's analyse a comparative governance account, helpful for designing ex ante oversight of algorithmic contracting in sensitive sectors.

¹⁰⁰ See, on baseline international standards, P Cihon, 'Standards for AI governance: international standards to enable global coordination in AI research & development' (2019) 40(3) *Future of Humanity Institute – University of Oxford* 340 ff.

¹⁰¹ See, for model laws and clauses, R Glowacki, 'Drafting AI Clauses: Five Tips' (2024) 7 *J Robotics Artif Intell & L* 301 ff. See, on cross-border redress mechanisms, O Pakuhinezhad and N Afsham, 'AI Regulation in Cross-Border E-Commerce: Legal Challenges and Global Remedies' (2025) 14(6) *Int J Sci Res* 44 ff.

for algorithmic fairness where consumers anywhere can raise issues about global platforms, and those platforms agree to abide by the ombuds recommendations (something analogous exists in some sectors, like the consumer financial ombuds in some multi-national banking contexts). Forth, Regulatory Cooperation Protocols: formalize the type of cooperation we see emerging. For instance, a memorandum of understanding between FTC (US) and European Commission on AI enforcement to share information on investigations, or a treaty that if a company is penalized in one jurisdiction for AI contract violations, that penalty can be enforced on the company's assets in another jurisdiction if the company doesn't pay (similar to how some cross-border financial penalties work).¹⁰² Fifth, Capacity Building and Norm Diffusion: not all countries have resources to develop AI rules from scratch. International organizations can help share best practices, toolkits, even software for auditing algorithms. If regulators around the world use similar tools (e.g., an open source algorithmic bias detection suite), they might converge on standards of what constitutes "unacceptable bias" numerically. Capacity building also helps avoid a scenario where AI companies simply go to places with lax oversight because those places don't have ability to monitor—the aim would be that every major market has at least a baseline of capability to audit and enforce AI rules.¹⁰³ Finally, Industry self-regulation as a bridge: sometimes industry groups create codes of conduct that go global. For example, something like a "Global AI Charter for Consumer Services" endorsed by big tech companies could fill gaps while law catches up. If major market players commit to, say, not using AI to target vulnerable consumers with unfair contract terms (and have an independent panel overseeing compliance), that could set an expectation that bleeds into legal norms. Eventually, those voluntary rules might become mandatory as regulators pick them up.¹⁰⁴

Nevertheless, cultural and legal differences persist (e.g., the U.S. First Amendment can affect how far you can regulate AI content in contracts or advertising), but as with privacy law, where a core set of principles (transparency, purpose limitation, data quality, etc.) are now common even if implemented variably, AI contract law could see a core set of guardrails universally recognized. The benefit will be twofold and it will consist in protecting consumers everywhere to a reasonable degree, and in giving businesses clarity and consistency to foster innovation that is also trustworthy.

The overarching aim is to create a market environment where doing the right thing with AI becomes a competitive advantage, not a handicap. If only the slowest, most cautious companies follow the rules and others cut corners for short-term gain, that is a problem. But if we reward ethical AI, then others will

¹⁰² See, on regulatory cooperation protocols, N Von Ingersleben-Seip, 'Competition and Cooperation in Artificial Intelligence Standard setting: Explaining Emergent Patterns' (2023) 40(5) Rev Pol'y Res 781 ff.

¹⁰³ See, on capacity building and norm diffusion, HF Hansen and others, 'Understanding Artificial Intelligence Diffusion Through an AI Capability Maturity Model' (2024) 26(6) Inf Syst Front 2147 ff.

¹⁰⁴ See, on Industry self-regulation as a bridge, T Marwala, *The Balancing Problem in the Governance of Artificial Intelligence* (Singapore 2024) 207 ff; DE Walters and HJ Wiseman, 'Self-regulation in Emerging and Innovative Industries' (2024) 62(3) Hous L Rev 543 ff.

emulate it to get the same benefits, and ideally, this could lead to a race-to-the-top dynamic in certain aspects of AI usage in contracts.

In any case, regulatory efforts can only go so far if companies themselves are not structured to handle AI issues.¹⁰⁵ Therefore, a key issue is to integrate AI governance into corporate compliance and risk management via the institutionalization of AI contract governance within corporate structures.¹⁰⁶

Institutionalizing these practices creates a first line of defense, and it means when regulators or journalists come knocking, the company is not scrambling since they have documentation of their processes and decisions. Moreover, it drives a message that AI responsibility is everyone's job in the company, reducing the likelihood of ethical corners being cut for expediency.

11. Integrating AI Regulation into Consumer and Contract Law Reform.

A durable reform strategy should mainstream artificial intelligence (AI) considerations into the interpretation and evolution of existing consumer protection and contract law, rather than treating "AI law" as an isolated silo. This integrationist posture is consistent with leading international governance instruments that frame AI oversight as a matter of aligning socio-technical systems with rule-of-law values, rights protection, and accountability, while remaining technology-neutral enough to survive rapid change.¹⁰⁷ It is also increasingly reflected in legally binding and quasi-binding frameworks that adopt risk-sensitive approaches and insist on transparency and responsibility across the AI lifecycle.¹⁰⁸ Against this background, the central private-law question is not whether AI is "special," but how established doctrines of formation, validity, fairness, consent, and responsibility should be applied (and, where necessary, refined) when contractual interactions are mediated or shaped by algorithmic systems.¹⁰⁹

First, legislatures and courts should clarify how core doctrines of contract formation operate in AI-personalised and dynamically priced environments. Classical common-law reasoning distinguishes offers from invitations to treat, and locates the moment of contractual commitment in the exchange of a legally operative offer and an acceptance that objectively corresponds to it.¹¹⁰

¹⁰⁵ I Rudko and others, 'New Institutional Theory and AI: Toward Rethinking of Artificial Intelligence in Organizations' (2025) 31(2) J Manag Hist 261–72.

¹⁰⁶ For example, appointing AI accountability officers, nominating AI governance committees, providing internal audits and compliance integration, incorporating ethical design and testing protocols, training and creating awareness in the workforce.

¹⁰⁷ OECD, 'Recommendation of the Council on Artificial Intelligence' (OECD/LEGAL/0449, 22 May 2019) <https://legalinstruments.oecd.org/en/instruments/oecd-legal-0449> accessed 30 June 2026.

¹⁰⁸ Council of Europe, 'Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law' (5 September 2024) <https://www.coe.int/en/web/artificial-intelligence/the-framework-convention-on-artificial-intelligence> accessed 30 June 2026.

¹⁰⁹ UNESCO, 'Recommendation on the Ethics of Artificial Intelligence' (2022) <https://www.unesco.org/en/articles/recommendation-ethics-artificial-intelligence> accessed 30 June 2026.

¹¹⁰ See generally the classic 'display/invitation to treat' line of authority as a foundation for classifying retail and platform communications in offer-and-acceptance analysis.

Where an online interface, recommender, or pricing algorithm generates multiple personalised variants, private law must determine whether each variant is an “offer” capable of acceptance or merely a solicitation to bargain. Established doctrine in retail display contexts treats many “displays with price” as invitations to treat rather than binding offers, partly to avoid compelling traders to supply in unlimited quantities on the basis of communications that are not intended to be final commitments.¹¹¹ Yet contract doctrine also recognises that some public-facing statements – especially those that are sufficiently definite and evince an intention to be bound – may constitute offers capable of acceptance.¹¹² AI-personalised contracting intensifies this boundary problem and warrants doctrinal guidance so that courts can classify algorithmic communications in a principled way, without giving firms an easy “automation exception” to ordinary formation rules.

Secondly, the doctrine of mistake should be modernised, or at least authoritatively re-articulated, for AI-driven errors, particularly pricing glitches and automated misrepresentations. The most difficult cases are those in which consumers “accept” an implausibly low price generated by a system error and the trader seeks to avoid the contract by invoking unilateral mistake. Courts addressing online pricing mistakes have treated the issue through orthodox principles, asking, among other things, whether a reasonable counterparty knew or ought to have known of the mistake and whether there was genuine consensus ad idem.¹¹³ The reform objective here is not to create a special AI mistake doctrine, but to ensure predictable allocation of risk between trader and consumer where automated systems produce misleading contractual signals. Private law should avoid outcomes that let firms externalise operational risk onto consumers (“the algorithm made me do it”), while still allowing avoidance in truly opportunistic “snap-up” scenarios where a consumer knowingly exploits an obvious error.¹¹⁴

Thirdly, consumer-consent standards should be tightened where AI personalisation affects price, core terms, or legally salient rights. Meaningful consent is strained when the consumer cannot see the baseline against which personalisation is occurring, or cannot reasonably understand why they received one set of terms rather than another. Contemporary consumer law is already moving toward requiring explicit disclosure of algorithmic personalisation in certain settings, including where the “price” presented has

¹¹¹ See *Fisher v Bell* 1 QB 394, where the display in shop window is treated as invitation to treat rather than an offer.

¹¹² See *Carlill v Carbolic Smoke Ball Co* 1 QB 256, where it is established that the advertisement can constitute an offer where sufficiently definite and intended to be binding.

¹¹³ See *Chwee Kin Keong and Others v Digilandmall.com Pte Ltd* SGHC 71 (High Court) and *Chwee Kin Keong and Others v Digilandmall.com Pte Ltd* SGCA 2 (Court of Appeal), where the online contracting and unilateral mistake in mispricing scenario is presented.

¹¹⁴ The ‘snap-up’ limitation on enforcing obvious unilateral mistakes is discussed in connection with the online mispricing context in the Digilandmall decisions. See E Mik, ‘The resilience of contract law in light of technological change’ in M Furmston (ed), *The Future of the Law of Contract* (Routledge 2020) 112 ff; U Tiwari, ‘Price Dominance Doctrine: Reclaiming the Doctrine of Offer in Modern Contract Law’ (2025) https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5970134 accessed 23 April 2026.

been personalised through automated decision-making.¹¹⁵ That logic should be extended in private-law reform by i) demanding clear notice that material contractual terms were personalised, ii) requiring explicit assent where personalisation materially reduces statutory rights (such as cooling-off periods, termination friction, or remedial protections), and iii) promoting periodic re-consent in long-running, AI-mediated contracts where the system iteratively re-optimises terms over time. These approaches align with broader data-protection constraints on solely automated decisions that produce legal or similarly significant effects, and with safeguards that emphasise contestability and human intervention.¹¹⁶ They also track the growing regulatory concern with manipulative interface design that undermines free and informed consumer choice.¹¹⁷ Private-law reform can borrow these insights to treat «consent obtained through algorithmic choice architecture» as potentially defective where the interface materially impairs autonomous decision-making or obscures the contractual reality.¹¹⁸

Fourthly, the private-law tests for unfairness and unconscionability should explicitly recognise algorithmic exploitation as a relevant contextual factor. Consumer-contract fairness review already permits courts to consider the circumstances surrounding contract conclusion, whether terms were pre-formulated, and whether a significant imbalance contrary to good faith exists.¹¹⁹ In AI contexts, a term may appear facially neutral yet be produced by a process that targets behavioural vulnerabilities, informational asymmetries, or predicted inertia, for example, using propensity models to impose harsher terms on consumers unlikely to complain or switch. Doctrinal guidance should make clear that such procedural unfairness, especially when combined with substantive detriment, can support findings of unfairness, deception, or aggressive practice, even if the term is not obviously abusive when read in isolation.¹²⁰ This is a straightforward “translation” of longstanding private-law equity concerns into the modern informational environment: the wrong lies not

¹¹⁵ Directive 2011/83/EU of the European Parliament and of the Council of 25 October 2011 on consumer rights, amending Council Directive 93/13/EEC and Directive 1999/44/EC of the European Parliament and of the Council and repealing Council Directive 85/577/EEC and Directive 97/7/EC of the European Parliament and of the Council OJ L304/64, art 6(1)(ea), as inserted by Directive (EU) 2019/2161.

¹¹⁶ Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation) [2016] OJ L119/1, art 22.

¹¹⁷ Regulation (EU) 2022/2065 of the European Parliament and of the Council of 19 October 2022 on a Single Market For Digital Services and amending Directive 2000/31/EC (Digital Services Act) OJ L277/1, art 25.

¹¹⁸ US FTC (n 32), surveying how manipulative choice architecture can undermine consumer autonomy and be treated as unfair or deceptive.

¹¹⁹ OECD, ‘Recommendation of the Council on Consumer Protection in E-commerce’ (12 May 2016) https://www.oecd.org/en/publications/oecd-recommendation-of-the-council-on-consumer-protection-in-e-commerce_9789264255258-en.html accessed 30 June 2026.

¹²⁰ Council Directive 93/13/EEC of 5 April 1993 on unfair terms in consumer contracts OJ L95/29, arts 3–5.

only in the clause, but also in the exploitative path by which the clause was produced and imposed.¹²¹

Fifthly, law reform should consider whether certain high-power, data-intensive intermediaries should bear heightened duties of loyalty and care in their AI-mediated dealings with consumers. Influential scholarship has proposed treating some digital actors as “information fiduciaries”, analogising their structural informational power to traditional fiduciary relationships in which the law restrains opportunism.¹²² Although fiduciary concepts are not historically central to mass consumer contracting,¹²³ recent legislative proposals have “flirted” with imposing duty-based governance on data-driven platforms, including duties of care, loyalty, and confidentiality in relation to personal data practices,¹²⁴ and it is worth noting that similar approaches appear in U.S. federal legislative proposals framed around “data care”.¹²⁵ Even if jurisdictions resist full fiduciary transplantation, private-law reform can still adopt the underlying logic: where AI systems create extreme information and power asymmetries, the law can justify heightened obligations to avoid manipulation, conflicts of interest, and harmful optimisation.

Sixthly, AI harms in consumer contracting can be intensified by weak competition; accordingly, competition law should be treated as a complementary private-law safeguard. Where algorithmic pricing and behavioural targeting converge across a market, consumer choice may cease to function as a meaningful constraint. Competition agencies and international bodies have explicitly analysed the risk that algorithms can facilitate coordination or parallel conduct, including in ways that challenge traditional evidentiary models of collusion.¹²⁶ Major enforcement communities have also examined how common algorithmic vendors, shared datasets, or standardised pricing tools can yield uniform market behaviour, raising questions about tacit coordination and market power.¹²⁷ In parallel, consumer law can be

¹²¹ Directive 2005/29/EC of the European Parliament and of the Council of 11 May 2005 concerning unfair business-to-consumer commercial practices in the internal market and amending Council Directive 84/450/EEC, Directives 97/7/EC, 98/27/EC and 2002/65/EC of the European Parliament and of the Council and Regulation (EC) No 2006/2004 of the European Parliament and of the Council (Unfair Commercial Practices Directive) OJ L149/22, where it is stated a general prohibition of unfair commercial practices, including misleading and aggressive practices; European Commission, ‘Guidance on the interpretation and application of Directive 2005/29/EC of the European Parliament and of the Council concerning unfair business-to-consumer commercial practices in the internal market’ OJ C526/1 [https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:52021XC1229\(05\)](https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:52021XC1229(05)) accessed 30 June 2026.

¹²² J.M. BALKIN, ‘Information Fiduciaries and the First Amendment’, (2016) 49 (4) Yale Law Journ., 1205 ff.

¹²³ See G Alpa and V Zeno-Zencovich, *Italian Private Law* (London 2007), where the A.’s develop a useful analysis for framing how general private-law categories are reinterpreted in mass consumer contracting and platform-mediated transactions.

¹²⁴ Balkin (n 120).

¹²⁵ See *New York Privacy Act proposal (A680)*, where data fiduciary framing, duties including care, loyalty, and confidentiality are treated.

¹²⁶ See U.S. ‘*Data Care*’ proposals, including the *Data Care Act framework* as legislative texts proposing duties such as care/loyalty/confidentiality in data processing contexts.

¹²⁷ OECD, ‘Algorithms and Collusion: Competition Policy in the Digital Age’ (2017) <https://www.oecd.org/content/dam/oecd/en/publications/reports/2017/05/algorithms-and->

strengthened by ensuring that exploitative AI contracting practices are not protected by market-wide normalisation. The policy point is structural: robust competition increases the practical value of consumer rights by providing credible exit options when one firm's AI-driven terms are oppressive.¹²⁸

Finally, coherent reform requires cross-disciplinary alignment rather than fragmented legal "patches" as AI in contracting touches privacy, consumer protection, platform governance, anti-discrimination, and safety. Soft-law tools that operationalise risk governance and accountability, while remaining adaptable, can support doctrinal integration by supplying shared vocabularies (risk mapping, impact assessment, documentation, monitoring) that courts and regulators can recognise across domains.¹²⁹ At the same time, hard-law AI regimes increasingly impose transparency expectations (for example, obligations to inform people when they are interacting with AI systems, and to disclose certain AI-generated or manipulated content), which can serve as interpretive "signals" for how private-law doctrines should treat concealment, misrepresentation, and defective consent in AI-mediated contracting.¹³⁰ National strategies that privilege regulator-led, adaptable frameworks similarly underscore the value of embedding AI considerations into existing legal institutions rather than building a monolithic, standalone "AI code" detached from everyday private-law adjudication.¹³¹

Therefore, the goal is to treat AI as a powerful modality through which businesses and consumers transact, and therefore as a factor that must be absorbed into the ordinary grammar of private law: formation, mistake, consent, fairness, and responsibility. Over time, as AI becomes ubiquitous, "AI-specific" rules should increasingly operate as refinements within mainstream consumer and contract doctrines, precisely because the legitimacy and stability of the legal system depend on continuity in core private-law principles, even as the technologies that mediate transactions evolve.¹³²

[collusion-competition-policy-in-the-digital-age_02371a73/258dcb14-en.pdf](#) accessed 30 June 2026, where it is conducted a competition-policy analysis of algorithmic coordination risks.

¹²⁸ Autorité de la concurrence, 'Algorithms and Competition joint study' (November 2019) <https://www.autoritedelaconcurrence.fr/sites/default/files/algorithms-and-competition.pdf> accessed 30 June 2026, where it is conducted an analysis of competition issues raised by algorithmic systems.

¹²⁹ Reference can be made, in this sense, to the materials on pricing algorithms/personalised pricing and competition/consumer risks in digital markets developed by the UK Competition and Markets Authority (CMA).

¹³⁰ National Institute of Standards and Technology (NIST), 'Artificial Intelligence Risk Management Framework (AI RMF 1.0)' (January 2023) <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf> accessed 30 June 2026, where it is traced a voluntary governance framework for managing AI risks and promoting trustworthy AI.

¹³¹ Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts OJ L2024/1689, art 50, which provides transparency obligations, including informing people they are interacting with an AI system and disclosure duties for certain AI-generated/manipulated content.

¹³² UK Government, 'A pro-innovation approach to AI regulation' (White Paper, 2023) <https://www.gov.uk/government/publications/ai-regulation-a-pro-innovation-approach/white-paper> accessed 30 June 2026, where regulator-led, adaptable framework, emphasis on interoperability and embedding AI governance into existing regulatory structures are considered very relevant.

In essence, the legal system should evolve to see AI as just another factor in how businesses and consumers interact, albeit a powerful one. That means mainstreaming AI considerations into the general doctrines of fairness, consent, and liability, rather than treating “AI law” as separate. Over time, as AI becomes ubiquitous, having a separate body of AI law might even fade, the special rules would merge into standard law. But right now, an explicit integration effort is needed to reach that mature state.

12. Toward Responsible AI in Consumer Contracts: Synthesis and Future Outlook.

The policy recommendations above collectively aim to create an ecosystem where AI can be safely and effectively harnessed for consumer benefit without eroding trust or rights. They span international coordination down to corporate practice, reflecting the multifaceted challenge at hand.

A central theme is balance via encouraging innovation and beneficial uses of AI (through sandboxes, incentives, harmonization to reduce red tape) while imposing boundaries to prevent abuse (through duties, oversight, and enforcement). The recommendations emphasize that protecting consumers and fostering innovation are not mutually exclusive, and, indeed, consumer trust is a prerequisite for sustainable innovation. If people fear AI or find themselves consistently harmed by it in marketplace dealings, they will resist its adoption and regulators will clamp down harder. Conversely, demonstrate an AI environment that respects rights and people will embrace its conveniences.

Another theme is shared responsibility since consumers, companies, regulators, and international bodies all have roles. While consumers empowered with rights and information can act as watchdogs and decision-makers in their own interest, companies with integrated governance and ethics can self-police much better than any external authority alone could achieve. At the same time regulators equipped with expertise and modern tools can act surgically and insightfully rather than bluntly, and international cooperation can prevent the cracks through which bad practices slip.

Implementing these recommendations will require political will, industry buy-in, and perhaps some trial and error. But the cost of inaction, such as eroding consumer protection, uneven playing fields, and stunted innovation due to lack of trust, justifies a proactive approach. Law often lags technology, but with AI in consumer contracts there is an opportunity to catch up and even get ahead of the curve by learning from the early signals we already see.

Ultimately, the measure of success will be when AI-powered consumer contracts become “normal” contracts in the eyes of the law, and that is, when their AI element is fully accounted for in legal analysis, neither a mystery nor a scapegoat.¹³³ Consumers should be able to engage in these contracts with no

¹³³ See F Martin-Bariteau and M Pavlovic, ‘AI and contract law’ in F Martin-Bariteau and T Scassa (eds), *Artificial Intelligence and the Law in Canada* (LexisNexis 2021) <https://ssrn.com/abstract=3730385> accessed 23 April 2026, where the A.’s underline that the growing use of artificial intelligence (AI) in contracting does not, in itself, undermine contract law,

greater caution or skepticism than they would a traditional contract, because the safeguards are in place. And businesses should be able to deploy AI confidently, knowing the rules of the road and that doing right by the customer is aligned with their legal obligations and profitability.

Artificial intelligence is no longer a future abstraction, but it is a present reality transforming the legal landscape of consumer contracts. From auto-generated terms in subscription services to algorithmic price discrimination and dynamic risk assessments in insurance and credit, AI systems now play a decisive role in how consumers engage with, consent to, and are bound by legal obligations. These systems offer tremendous efficiency gains and enable products and services tailored to individual needs, but they also raise unprecedented legal risks that challenge the foundations of contract and consumer protection law.

In concluding, this article posits that the way forward lies in a rights-based, consumer-centric framework that is also pragmatic about technological realities. It shall be demanded that AI respect and empower consumers, not undermine them, hence transparency, fairness, and accountability are non-negotiable. At the same time, regulation should be sufficiently nuanced and tech-informed so as not to smother beneficial innovations or entrench incumbent players (who can afford compliance costs) at the expense of new entrants.

The era of AI-powered consumer contracts is still in its early chapters. As lawmakers, regulators, industry leaders, and academics continue to write the rules of this new game, the guiding vision should be one of “responsible automation”: harnessing AI’s capabilities to enhance consumer welfare (through better products, lower costs, personalized services) while diligently curbing its potential to mislead, discriminate, or exploit. Achieving this will not be easy—it demands interdisciplinary effort, continuous oversight, and willingness to adjust course as we learn from real-world deployments. But the payoff is a digital marketplace where innovation and consumer trust grow in tandem, rather than in tension.

In sum, regulating AI-powered consumer contracts is not about erecting barriers to technology, but it is about erecting guardrails that ensure technology serves the public interest. With adaptive and principled legal responses, it can be welcomed the age of algorithmic contracting not as a threat to consumer rights, but as an opportunity to reaffirm and strengthen those rights in a transformative era.

which is largely capable of accommodating contracts drafted, negotiated, concluded, or performed through automated systems. Adopting a functional approach, both civil law and common law doctrines can interact effectively with AI-mediated contractual practices. At the same time, the use of AI in drafting, forming, or enforcing contracts introduces significant risks for controllers, legal professionals, and especially vulnerable parties. While these concerns are not new, AI can intensify existing power imbalances and exacerbate social injustices, and this development highlights the need to consider the harmonisation and adaptation of the current contractual framework to AI-driven contracts, with particular attention to transparency, accountability, and the protection of weaker parties.

