

Better Markets Response to the House Financial Services Committee Democrats' Request for Information on Artificial Intelligence in the Financial Marketplace

August 14, 2026

The Hon. Maxine Waters
Committee on Financial Services
United States House of Representatives
2129 Rayburn House Office Building
Washington, DC 20515

Dear Ranking Member Waters:

On behalf of Better Markets, thank you for the opportunity to respond to the Committee Democrats' Request for Information on artificial intelligence (AI) use in the financial marketplace. Better Markets is a non-partisan, public-interest 501(c)(3) nonprofit that has fought for 15 years for a financial system that serves Main Street workers, families, small businesses, and communities, and not just Wall Street. We have brought that same fight to the questions AI now poses.

The accompanying submission responds to the questions on which our institutional expertise and published work allow us to add the most value, among them model fairness and explainability, model risk management, data privacy and disclosures, third-party service providers and the competitive position of community banks and community development financial institutions (CDFIs), liability, capital markets and investor protection, housing, workforce, data centers, and financial stability. Where questions fall outside the scope of our work, we have not responded; our silence on those questions should not be read as agreement or disagreement with any premise they contain.

We commend the Committee for undertaking this inquiry while the choices are still open. Our animating conviction is a simple one: change is certain, but progress is not. Whether AI makes economic life better or worse for most Americans depends on the rules written around it, and we welcome the chance to work with you and the Committee to ensure the public interest has a seat at the table before those rules harden.

Thank you for the opportunity to provide these comments. We look forward to continuing to engage with the Committee as its work progresses.

Sincerely,

[Evan R. LeFlore](#) | *Director of AI, Innovation, and Economic Opportunity* | Better Markets, Inc.



Preliminary Submission Details

Question #1: Name and organization.

This response is submitted by Better Markets, Inc., a non-partisan, public-interest 501(c)(3) non-profit based in Washington, D.C. For 15 years, Better Markets has advocated for greater transparency, accountability, and oversight in the domestic and global capital and commodity markets to protect the economic security, opportunity, and prosperity of the American people. Our AI program is led by Evan LeFlore, the organization's inaugural Director of AI, Innovation, and Economic Opportunity.¹

Question #2: Supplemental materials.

Better Markets is submitting supplemental materials to AIRFI_Data@mail.house.gov, including our June 2026 report *AI, the Economy, and You: A People-Centered AI Agenda* and related fact sheets and analyses referenced throughout this response.

Question #3: Contact email.

The Committee may contact us for further details at eleflore@bettermarkets.org or (202) 618-6414.

Executive Summary

Artificial intelligence is not primarily a technology story. It is an economic story about who benefits, who pays, and who is accountable when automated systems make the most consequential financial decisions of people's lives, including whether they get a loan, a card, an insurance payout, or a job. The people building the technology are warning, in unusually blunt terms, that the disruption ahead will be severe and that its gains will not be broadly shared.² Congress has a narrow window to ensure that the public interest has a seat at the table before the rules harden.

These stakes are not distributed evenly. Because AI systems learn from historical data that encodes decades of redlining, discriminatory lending, and unequal access to credit and housing, they risk reproducing those patterns at machine speed and scale, while cloaking them in the appearance of neutral, objective math. The communities that bear the greatest risk are the same ones that prior waves of financial practice already disadvantaged: Black, Hispanic, and other borrowers of color, the mission-driven institutions that serve them, and young workers of color entering the labor market. Whether AI becomes a tool that widens access or one that automates exclusion depends on the rules written around it. Throughout this RFI response, we flag where AI's risks fall disproportionately on communities of color, and where the safeguards we recommend are also the safeguards that protect them.

The RFI arrives at a moment when the federal government is moving in the opposite direction from the one this time demands. The Consumer Financial Protection Bureau (CFPB) has been subjected to a sustained campaign to defund it, fire its staff, halt its supervision, and rescind its guidance, including its key guidance on adverse action notices and AI-driven credit decisioning.³ At the same time, banking regulators' April 2026 revised model risk management guidance (SR 26-2) expressly carved generative and agentic AI out of its scope at precisely when the technology is being given the ability to make decisions that consequentially impact Americans' lives and livelihoods.⁴ The result is a widening accountability gap, as AI is being deployed the fastest exactly where oversight is being withdrawn fastest.

Our recommendations cluster around six themes:

- **Explainability is a legal obligation, not an aspiration, and disparate impact review must be protected.** If a lender cannot explain why an AI model denied someone credit, the model is not ready for deployment. The Equal Credit Opportunity Act's (ECOA) adverse-action requirements already demand specific, accurate reasons; regulators should enforce them and reject the claim that model complexity excuses non-compliance.⁵ Because AI models discriminate largely through the use of more "neutral" proxy data, disparate-impact analysis is often the only way to detect that discrimination. Congress should codify this before it is further stripped away by executive action.
- **Model risk management must cover the models that actually pose risks.** Excluding agentic and generative AI from SR 26-2 leaves banks to self-govern the highest-risk, least-understood systems with no rulebook. That carve-out should be reversed.

- **Liability must be clear before these systems become entrenched.** When an autonomous system produces a discriminatory or harmful outcome at scale, someone must be held legally accountable. Ambiguity today guarantees a larger crisis and reactive over-regulation tomorrow.
- **Concentration itself is a risk.** The concentration of a few dominant third-party AI models threatens both competition (community banks and CDFIs cannot replicate larger banks' data advantages) and financial stability (through correlated, herd-like behavior amplifying shocks).
- **The costs of the buildout are being shifted to the public.** Data centers are driving regressive increases in electricity bills and imposing environmental burdens on communities that see little of the benefit. Those costs should be borne by those capturing the gains.
- **Existing law and existing authorities go a long way, if regulators want to use them.** Much of what is needed does not require new statutes. It does require that agencies apply anti-discrimination, consumer-protection, securities, and safety-and-soundness laws to AI. It also requires that Congress preserve the enforcement and supervisory capacity to do so. Where the new processes are created, Congress should guard against the new risks they introduce.

The through-line is that AI's risks in finance are not hypothetical or distant. Discriminatory AI underwriting has already produced real settlements and audit findings.⁶ Algorithmic trading already dominates markets and has already produced flash-crash cascades with no accountable party.⁷ And AI credit models are already producing denials that the institutions deploying them cannot explain as the law requires. The task before the Committee is to close these gaps while the choices are still open.

White House Actions on AI

Question #4: Federal review process for covered frontier models.

Better Markets does not take a position on the technical design of frontier-model review process established by the June 2, 2026 executive order. We offer two principles drawn directly from our work.

First, any review process must be genuinely independent of the firms being reviewed and must be resourced to match the pace and complexity of the technology. Our foundational concern with the current trajectory of AI policy is that private corporations racing to be first and choosing short-term profits are driving AI's development while the public is being viewed as an impediment to innovation. A voluntary review overseen by an under-resourced agency, or one captured by the perspectives of the developers it reviews, would reproduce that dynamic rather than correct it. If



Treasury is to have a role, its participation should be structured to embed a Main Street perspective, including the interests of workers, consumers, depositors, and small businesses.

Second, the Committee should be alert to a distinct market-integrity concern around the review process, which gives the government a 30-day advance look at unreleased frontier models. Such concerns include the trajectory, capabilities, and commercial prospects of publicly traded AI developer companies and their partners. Government personnel, contractors, and reviewers who obtain early knowledge of a model's capabilities (or its failure to clear review) would be privy to valuable information positioning them to trade in securities of the developers themselves and the many public companies whose valuations now move on AI expectations. The same asymmetry could invite front running of the market reaction to a model's release or rejection. Before the process is operationalized, Congress and Treasury should ensure that robust information barriers, trading restrictions, and disclosure and enforcement mechanisms cover everyone with access, so that a public-interest review does not become a vector for insider trading.

Model Fairness and Explainability

Question #10: Explainability standards for AI decisions such as credit denials.

This is among the most important questions in the RFI, and Better Markets' position is clear: if an AI system used in underwriting cannot produce a specific and accurate reason for an adverse action, it should not be used for that purpose.

The law already requires this. Under ECOA and Regulation B, a creditor taking adverse action must provide the specific, principal reasons for the denial. The CFPB's now-rescinded September 2023 guidance correctly explained that lenders "cannot simply use CFPB sample adverse action forms and checklists if they do not reflect the actual reason for a denial of credit."⁸ The rescission of that guidance did not change the statute, and ECOA's requirements remain in force. The Committee should insist that regulators enforce them against AI models exactly as against any other underwriting method. The harm caused by this guidance rescission is compounded by the CFPB's amendments to Regulation B, which seek to remove disparate impact as an avenue to establish legal liability for credit decisions made by lenders.

The urgency is not theoretical. The CFPB's own Winter 2025 Supervisory Highlights found that AI-driven credit scoring models at banks were producing loan denials that the institutions themselves could not explain as required by law.⁹ When a machine-learning model with thousands of interacting variables denies a loan, and even its creators cannot trace the decision to an identifiable cause, the borrower's statutory right to a real explanation is dead on arrival.

This failure does not land randomly. When a model's reasoning is opaque, the applicants least able to absorb an unexplained denial (and least likely to have the resources to contest it) are disproportionately borrowers of color. An explainability requirement is therefore not a technical nicety, it is a civil-rights safeguard.

We therefore recommend that Congress and regulators establish, at minimum:

- **An enforceable explainability floor.** Any AI system used in a decision covered by ECOA, the Fair Housing Act, or the Fair Credit Reporting Act must be capable of generating specific, accurate, and individualized reasons for adverse action. The burden of demonstrating that capability should rest on the deploying institution before deployment, not on the harmed consumer after the fact.
- **A prohibition on treating model complexity as a compliance defense.** “The model is a black box” is a reason not to deploy the model for a high-stakes decision, not an excuse for failing to explain the decision.
- **Meaningful notices to affected individuals that AI was materially involved in the decision.** This should be taken together with the actual reasons and a workable path to contest an error.
- **Restore disparate impact liability** by reversing recent changes to Regulation B adopted earlier this year.

Question #11: Pre-deployment and post-deployment controls.

The central lesson of the cases already on record is that discrimination frequently enters through inputs and proxy data, and that it is detectable only through disciplined pre- and post-deployment testing. In July 2025, the Massachusetts Attorney General settled with student-loan company Earnest Operations for \$2.5 million after finding that its AI underwriting model produced discriminatory outcomes for Black and Hispanic applicants by using cohort default rates from applicants' colleges.¹⁰ This data showed that the company disproportionately penalized graduates of Historically Black Colleges and Universities, including applicants with strong credit scores. In addition, the independent fair-lending monitorship of Upstart's model, conducted by Relman Colfax, found approval-rate disparities for Black applicants that were "statistically and practically significant," and the company declined to adopt an available less-discriminatory alternative model.¹¹

These cases point to the need for concrete controls:

- **In pre-deployment,** produce disparate-impact testing across protected classes, scrutinize every input and proxy variable for correlation with protected characteristics, document a mandatory search for less-discriminatory alternatives with a duty to adopt a materially less discriminatory alternative, and validate and document the model's data provenance and assumptions.
- **In post-deployment,** monitor outcomes for disparities. Also perform independent audits, including retention of the inputs, versions, and decision records needed to reconstruct any individual decision. Also create a governance obligation connected to a named, accountable human owner within the organization.



The unifying principle is that the institution deploying the model is responsible, not the consumer or algorithm developer. The institution bears the burden of proving a system is fair, explainable, and non-discriminatory on a continuing basis.

Finally, we urge the Committee to protect the most important tool for detecting these harms. Disparate-impact analysis is what surfaced the discrimination in both the Earnest and Upstart cases, both of which had statistically significant disparities that no neutral data point revealed on its own. Yet the Administration has directed agencies to "eliminate the use of disparate-impact liability in all contexts to the maximum degree," a move that, if carried into financial and housing supervision, would blind regulators to exactly the algorithmic discrimination this section is meant to catch.¹²

The danger is acute for AI, because machine-learning models can reproduce historical discrimination through neutral proxies such as college attended, ZIP code, or spending pattern. Disparate-impact review is often the only way to detect that a model is discriminating.

Better Markets therefore urges the Committee to recognize the essential role of disparate-impact review in the AI era and to codify it. This places the disparate-impact framework under ECOA and the Fair Housing Act on a firm statutory footing so it cannot be rescinded by executive action. We return to this point in our response to Question 34.

Model Risk Management

Question #12: Changes to the Revised Guidance on Model Risk Management.

Better Markets recommends that the AI carve outs in the updated model risk management guidance be closed. When the Federal Reserve, Office of the Comptroller of the Currency (OCC), and Federal Deposit Insurance Corporation (FDIC) issued SR 26-2 in April 2026, they narrowed the definition of "model" and explicitly excluded generative and agentic AI from the guidance's scope, characterizing those technologies as "novel and rapidly evolving" and promising a separate request for information at a later date.¹³ The practical effect is that the framework governs the well-understood statistical models that pose comparatively modest and familiar risks, while exempting the novel, opaque, and rapidly proliferating systems that pose the greatest and least-understood risks. Banks are left to self-govern generative and agentic AI risk with no formal rulebook at precisely the moment these systems are making consequential decisions about people's financial lives.

We recommend that regulators:

- **Reverse the exclusion and bring generative and agentic AI within model-risk-management expectations.** This should be accomplished by tailoring the application of those expectations to the technology rather than exempting the technology from the framework altogether.

- **Align with model risk management expectations.** Regulators should recognize that the qualities regulators cited to justify the exclusion (i.e., novelty, rapid evolution, and complexity) are exactly the qualities SR 11-7 originally identified as increasing model risk and should continually be embedded into updated model risk guidance such as SR 26-2.¹⁴
- **Move with urgency.** Promising a future RFI while these systems scale is the kind of regulatory abdication that only guarantees bigger problems later. Policymakers should not repeat the pattern in which the government's response to a technological transformation came too late and did too little.

Question #13: Amending federal data-privacy law for AI while preserving identity verification and financial-crime compliance.

AI systems, and generative AI in particular, rely on collecting and using massive amounts of data to function, and may even include users' sensitive personal data. The United States lacks a comprehensive federal privacy framework, and existing protections are patchwork. Better Markets supports strengthening consumer control over personal financial data and cautions that any reform be written so that legitimate, tightly scoped uses are preserved, such as customer identification, fraud detection, and anti-money-laundering compliance. While foreclosing the repurposing of that same data for AI-driven surveillance, profiling, and consumer marketing must never be meaningfully authorized.

In general, Better Markets advocates that any federal framework related to AI be treated as a floor, not a ceiling, and this is also a guideline we hope Congress incorporates into any laws around data privacy.

The animating concern from our agenda is the risk that AI will turn personal data into a tool for surveillance and exploitation. Data minimization, purpose limitation, and a clear separation between compliance uses and commercial-exploitation uses are the principles that reconcile the two objectives in the Committee's question.

Question #15: Customer notice and disclosure regimes for AI use.

Consumers are increasingly subject to consequential decisions made by AI systems often without much testing or even awareness by the users. Better Markets supports a meaningful notice-and-disclosure regime under which a consumer is told, in plain language, when AI is materially used to make or substantially inform a decision about a product or service they receive. This should be true whether built in-house or supplied by a third-party vendor. For adverse-action decisions, consumers should be given the actual reasons and a route to contest them. (See our response to Question 10).

Disclosure obligations should attach to the institution that owns the customer relationship and cannot be evaded by outsourcing the decision to a vendor. At the end of the day, the deploying institution remains responsible for what it deploys.

Question #18: Standardizing how publicly traded companies report their use of AI.

Better Markets supports standardized, decision-useful disclosure of material AI use by public companies. Investors today cannot reliably compare how companies deploy AI, what risks that use creates, or whether a company's public AI claims are accurate.

This gap has already produced a recognizable form of securities fraud called "AI washing." This is the practice of making false, exaggerated, or misleading statements about the extent or sophistication of a company's use of artificial intelligence—for example, claiming a product is "AI-powered" when it is not, or overstating AI's role in generating investment returns—in order to attract investors, customers, or capital by riding the market's enthusiasm for AI. It is the AI-era analogue of "greenwashing."

This matters to the Committee because AI washing directly harms investors and distorts capital formation as it induces people to buy securities on the basis of capabilities that do not exist, misallocates capital toward firms that market AI rather than firms that actually deliver it, and erodes trust in the disclosures on which efficient markets depend. This is not hypothetical, as the Securities and Exchange Commission (SEC) has already charged investment advisers with exactly this conduct, making false and misleading statements about their use of AI.¹⁵

Standardized, specific, and comparable reporting, covering material AI use, associated risks, and governance, would give investors the information they need to separate substance from marketing and would narrow the space in which AI washing can occur. We note the tension between this need and the SEC's recent proposal to reduce the frequency of public-company disclosures, which in our view moves in the wrong direction for investor protection.¹⁶

See also our responses to Questions 23 and 31 on AI in issuer disclosures and substantiation of AI performance claims.

Third-Party Service Providers

Question #19: Heightening monitoring of emerging risks from AI in financial services and housing.

The concentration of AI capability among a small number of third parties is both a competition problem and a financial-stability problem. In the Committee's own bipartisan working-group sessions, the Federal Reserve raised concerns about "the problematic nature of having a 'monoculture of models,' whereby financial institutions all use the same third-party providers," and a panelist warned that widespread adoption of the same models may encourage herd-like behavior in capital markets.¹⁷ When many institutions rely on similar AI systems trained on similar data, they can produce correlated behavior such as buying and selling in lockstep that amplifies shocks rather than absorbing them.

We recommend that regulators use existing authorities to (a) map and monitor the concentration of critical third-party AI dependencies across the system, (b) treat significant common dependencies as a systemic-risk vector deserving of Financial Stability Oversight Council (FSOC)

and Office of Financial Research (OFR) attention (see Questions 49–51), and (c) require deploying institutions to retain full accountability for vendor-supplied models, including the testing, explainability, and monitoring obligations described above.

Question #20: Ensuring adequate oversight of third-party vendors during reduced enforcement and supervision.

This question hits at the heart of Better Markets’ concerns about this era of AI development. The most sophisticated AI-driven consumer finance activity is increasingly being conducted by non-banks with minimal regulatory oversight. This activity is scaling at the very moment where federal supervision and enforcement capacity is being deliberately reduced, particularly at the CFPB, where the Administration has sought to fire nearly all staff, halt supervision, drop enforcement actions, and rescind guidance. Before this era, the CFPB had returned more than \$21 billion to over 200 million consumers since 2011.¹⁸ Additionally, FSOC is currently in the process of scaling-back systemically important nonbank financial institution (SIFI) designation authority, which, if adopted, will make designation of nonbank SIFIs for heightened supervision all but impossible.

While Congress cannot substitute for the daily supervisory capacity of the CFPB, it can:

- **Preserve the enforcement and supervisory capacity of the agencies**, above all the CFPB, and resist efforts to defund, fold, or hollow them out. An on-paper regulator cannot oversee AI that it lacks the staff to examine.
- **Close the third-party-vendor oversight gap.** The problem this solves is a structural blind spot in financial supervision. When a bank or credit union outsources a critical function such as an AI underwriting, fraud-detection, or customer-service system to an outside technology vendor, the risk does not disappear, it simply moves to a company the regulator may have no authority to examine. Under the Bank Service Company Act, the banking agencies (Federal Reserve, OCC, FDIC) can examine third-party service providers as if the provider were the regulated institution itself following the risk to its source. But two important regulators lack that reach. The National Credit Union Administration (NCUA) has no statutory authority to examine the third-party technology vendors that serve credit unions, even though those credit unions increasingly depend on the same outside AI systems as banks; the NCUA has asked Congress for this authority for years. The SEC and Federal Housing Finance Agency (FHFA) both face an analogous gap in supervising the technology and AI vendors embedded in the finance ecosystems they oversee. The consequence is that a growing share of consequential AI in consumer finance and housing is operated by vendors that sit outside the examination perimeter, precisely the "monoculture of models" and concentration risk described in our response to Question 19 but beyond the reach of the regulators nominally responsible for the institutions that depend on them. Congress should give the NCUA and FHFA examination authority over third-party technology service providers comparable to what the banking agencies already have, a step FSOC and the GAO have repeatedly recommended and that has bipartisan legislative history.

- **Use its oversight powers** to insist that reduced agency capacity is not treated as a green light for supervised institutions to deploy unmonitored AI.

Question #21: Keeping community banks, CDFIs, and MDIs competitive in accessing AI.

Better Markets has long emphasized the central role of community banks and mission-driven lenders, including CDFIs and minority depository institutions (MDIs), to the economic health of local communities.¹⁹ The concentration of AI capability in the largest institutions is accelerating the consolidation of financial power, because community banks cannot match the data advantages and analytical capabilities of megabanks and nonbank lenders. Left to the market, most trend lines point toward further consolidation.

The stakes are especially high for MDIs and mission-driven CDFIs, which exist specifically to serve Black, Hispanic, Native American, and other communities that mainstream finance has historically underserved. If access to competitive AI capability concentrates among the largest institutions, these mission-driven lenders—who are already operating with thinner margins and smaller data sets—risk being pushed further to the margins, widening the very access gaps they were created to close. An AI transition that hollows out MDIs and CDFIs would not be a neutral market outcome, it would fall directly on the communities those institutions serve.

But this outcome is not inevitable. AI could actually help community banks, CDFIs, and MDIs level the playing field and become more efficient and competitive. We therefore recommend that Congress consider affirmative measures to prevent AI from becoming another engine of consolidation, including support for shared or cooperative AI infrastructure that smaller institutions can access on fair terms, attention to whether dominant vendors offer smaller institutions non-discriminatory pricing and access, and assurance that any new AI compliance obligations are calibrated so they do not fall with disproportionate weight on the smaller institutions least able to absorb fixed compliance costs without weakening the substantive protections consumers are owed regardless of their lender's size.

- **Treat the preservation of community banks, CDFIs, and MDIs as an explicit policy objective in AI infrastructure decisions.** Shared or cooperative AI infrastructure, fair vendor pricing, and calibrated compliance obligations should be designed with the survival and competitiveness of mission-driven lenders specifically in mind, given their outsized role in extending credit to communities of color.

Liability

Question #22: Liability for AI-driven violations of law that harm consumers.

Better Markets' position is a matter of principle: when an AI system makes a consequential financial decision and gets it wrong, the humans and institutions that built, deployed, and profited from that system must be accountable. Making liability clear before these systems become entrenched is key to a safe and well-functioning society.

The core legal problem is a chain-of-liability gap. For instance, when an autonomous system errs systematically across thousands of decisions, current law does not clearly say whether responsibility rests with the institution that deployed it, the vendor that built it, or the developer that trained the model. That ambiguity lets each point to the others while harmed consumers go without remedy.

The gap is not hypothetical. Denials by an autonomous system can occur at a scale and opacity that make individual redress illusory even at high error rates.²⁰ Algorithmic trading has already produced a cascading failure in the October 2025 crypto flash crash, when roughly \$19.3 billion in liquidations occurred within a day for which no single entity bore responsibility.²¹ And in the case of *Moffatt v. Air Canada*, a tribunal held a company fully liable for its customer-service chatbot's errors, rejecting the defense that the AI was a separate legal entity.²²

Meanwhile, existing anti-discrimination law already supplies the standard. For example, a human loan officer who systematically denied qualified Black applicants would face ECOA liability.²³ Yet an algorithm producing the same outcome through proxy variables sits in a legal gray zone regulators are only beginning to address.

Our recommendations follow directly from this diagnosis:

- **Establish clear liability that runs to the deploying institution for AI-driven violations of federal law that harm consumers** so that outsourcing a decision to an algorithm or to a third-party vendor does not launder away legal responsibility. Existing statutes (ECOA, the Fair Housing Act, Fair Credit Reporting Act (FCRA), Truth in Lending Act (TILA), Electronic Fund Transfer Act (EFTA), and the securities laws) supply the substantive standards; the task is to make unmistakably clear that they apply to AI-originated conduct and that liability cannot be evaded through automation. The outcome of the *Air Canada* case referenced above should be the baseline, not the exception.
- **Allocate responsibility across the chain rather than letting it fall through the cracks.** Deployer-facing liability should be the front door for the harmed consumer, but it should be paired with clear rules—through contract, regulation, or statute—that let deployers hold vendors and model developers accountable for defects they introduced so that responsibility ultimately rests with the party best positioned to prevent the harm. The goal is to eliminate the outcome in which no single entity bears responsibility.

- **Apply the fiduciary and consumer-protection standards we expect of human decision-makers to the AI systems replacing them.** There is no principled reason a duty owed by a human advisor, loan officer, or claims adjuster should evaporate when an autonomous system performs the same function. Better Markets has proposed developing fiduciary liability frameworks that hold AI systems to the same standards as human decision-makers (i.e., an "agentic fiduciary" standard) as a central plank of its agenda.²⁴
- **Approach safe harbors with great caution.** Better Markets does not categorically oppose narrowly drawn safe harbors tied to demonstrated, audited, good-faith compliance (for example, documented less-discriminatory-alternative testing). But a broad safe harbor that immunizes institutions from the consequences of the systems they deploy would be an invitation to deploy recklessly and would defeat the purpose of a liability framework. Any safe harbor must be earned through verifiable conduct, not granted by category.
- **Pair liability with transparency.** Liability is only enforceable if harms are detectable. Institutions deploying consequential AI should be required to preserve the records needed to reconstruct individual decisions and to support the disparate-impact review discussed in our responses to Questions 11 and 34, so that the accountable party can actually be identified after the fact.
- **Act now to avoid pendulum policymaking.** The industry's own self-interest should compel collaboration, because otherwise there will inevitably be a disaster that will result in a popular outcry and over-regulation.

Capital Markets, Investor Protection, and Capital Formation

Question #23: Standards for AI use by issuers in disclosures, periodic reports, and offering documents.

AI is increasingly used to generate the very disclosures on which investors and the officers who certify them rely. Better Markets supports clear standards making explicit that a company's existing obligations, including the accuracy of its disclosures, the integrity of its internal control, the officer certifications under Sarbanes-Oxley Sections 302 and 906, and liability under Sections 11 and 12 of the Securities Act and Rule 10b-5. We encourage these standards to be applied undiminished when AI is used to prepare those materials. AI use should be accompanied by human review and control adequate to the certifications the law requires. The use of an AI tool neither excuses an inaccurate disclosure nor dilutes an officer's responsibility for it. This concern is heightened by the SEC's recent proposal to reduce the frequency of public-company disclosures, which would give investors less of the material information they need at exactly the moment AI is complicating its reliability.²⁵

Question #24: Risks AI pose to market integrity.

The market-integrity risks are already materializing. Algorithmic trading accounts for roughly 60–70% of all trades, and an October 2025 crypto flash crash triggered approximately \$19.3 billion in forced liquidations and wiped out roughly 1.6 million accounts within a day after simultaneous algorithmic selloffs with no single entity bearing responsibility for the cascade.²⁶ Oxford researchers have documented that agentic AI in financial simulations exhibits behaviors no one programmed—including price collusion and market manipulation—that traditional oversight is not designed to detect.²⁷

We recommend that the SEC, Commodities Futures Trading Commission, FINRA, National Futures Association, and Municipal Securities Rulemaking Board each be equipped and directed to (a) modernize surveillance to detect coordinated or emergent manipulation by autonomous agents across platforms, (b) establish accountability for algorithmic cascades so that “no single entity bore responsibility” is no longer an acceptable outcome, and (c) treat synthetic-media-driven manipulation as squarely within the antifraud provisions of the securities laws, with the tools and staffing to prosecute it.

Question #25: Obligations for investment advisers and broker-dealers using AI in recommendations.

The fiduciary duty under the Investment Advisers Act, Regulation Best Interest, and the associated supervision and recordkeeping obligations should apply fully when an AI system, rather than a human, originates or shapes a recommendation. Better Markets is particularly concerned about conflicts of interest, including those introduced by reliance on third-party AI vendors, and about the adequacy of model testing, validation, and supervision. The core principle is continuity of duty: an adviser cannot discharge a fiduciary obligation by delegating the recommendation to an unvalidated, unexplainable, or conflicted AI system.

Question #29: Evidence on AI reducing intermediation costs and ensuring savings reach investors.

Better Markets supports the goal of ensuring that AI-driven efficiencies translate into lower costs for retail and institutional investors rather than being retained as additional firm revenue. This is a direct application of our central question of whether AI's gains are broadly shared or concentrated at the top. We caution that cost-savings claims should be substantiated rather than assumed, and that disclosure and, where appropriate, examination should test whether advertised efficiencies are in fact passed through to investors. The relevant metric is the net outcome for the saver, not the gross efficiency captured by the firm.

Question #31: Distinguishing substantiated AI performance improvements from marketing claims.

The SEC's “AI washing” enforcement actions demonstrate that unsubstantiated AI performance claims are already a live investor-protection problem.²⁸ Better Markets supports a clear substantiation standard where any AI-related performance representation made to investors

should be supported by evidence adequate to the claim, with the same antifraud exposure as any other performance representation. The burden should rest on the firm making the claim to substantiate it, not on the investor to disprove it.

Housing

Question #33: Transparency and clarity in PropTech algorithms.

The transparency and explainability principles in our responses to Questions 10 and 11 apply with full force to property technology (PropTech). PropTech tools utilize algorithms which are often trained on data that are rooted in discriminatory historical practices and which lack transparency, risking the further embedding of discriminatory patterns and disparate outcomes in the housing market.²⁹ The risk here, algorithmic redlining, has a history: automated valuation models that systematically undervalue homes in Black neighborhoods, and tenant-screening tools that convert past housing discrimination into present-day exclusion. Because homeownership remains the principal engine of wealth-building and the largest single source of the racial wealth gap, discriminatory AI in housing does not just deny individual transactions but instead entrenches a wealth divide that fair-housing law was enacted to dismantle.

Automated valuation models, automated underwriting systems, and tenant-screening tools should be subject to the same disparate-impact testing, input scrutiny, explainability, and auditability requirements we recommend for credit models generally, so that regulators and Congress can promote responsible innovation while ensuring access to fair and affordable housing and prohibiting discriminatory outcomes.

Question #34: FHFA and peer-agency oversight of PropTech and fair-lending compliance.

The Fair Housing Act, ECOA, and FCRA already prohibit the discriminatory outcomes PropTech can produce. This concern is sharpened by the Administration's rescission of key fair-lending regulations and its directive to agencies to eliminate the use of disparate-impact liability, which is a move that, if applied to AI-driven housing decisions, would remove one of the most important tools for detecting exactly the kind of statistically significant disparities uncovered in cases like *Earnest and Upstart*.³⁰ We recommend that the FHFA direct Fannie Mae and Freddie Mac to demonstrate and supervise fair-lending compliance for the PropTech and automated systems in their ecosystems, and that Congress use its oversight authority to ensure disparate-impact analysis remains available as a compliance and enforcement tool in housing. Preserving disparate-impact analysis is especially consequential in housing, where it has historically been the primary legal instrument for exposing lending and valuation practices that disadvantage communities of color without announcing that they do so.

Workforce and Training

Question #41: Addressing workforce challenges and job-displacement inequality.

Workforce displacement is a core pillar of Better Markets' AI agenda, and getting the diagnosis right is the precondition for getting the policy right. The best current evidence supports a single organizing finding: AI automates tasks, not whole jobs (at least so far) and the harm is landing unevenly on one group first. A U.S. Census survey of more than 100,000 firms found that augmentation of worker tasks is roughly four times more common than displacement, and that most AI-using firms report no employment change attributable to AI to date.³¹ But “tasks, not jobs” does not mean no one gets hurt, it means the hurt is concentrated and it points to sharper policy instruments than the mass-unemployment tools (broad stimulus, giant untargeted retraining) that a “jobs” framing would suggest. Get the diagnosis wrong, and every prescription is wrong too.

Two features of the evidence shape our recommendations. First, whether AI augments a worker or displaces one is decided not by the technology but by how firms choose to deploy it, which is precisely why policy has a seat at the table and why the augment-or-displace choice can be tilted by incentives. Second, the aggregate calm is not reassuring but precarious: the same task automation that registers as harmless “augmentation” today can curdle into displacement one budget cycle later. The uncertainty itself is the argument for acting now, while the outcome is still shapeable.

Better Markets therefore recommends a task-focused jobs agenda that prepares for both continued augmentation and the possibility of wholesale displacement:

- **Protect the bottom rung.** Entry-level jobs are where workers learn the judgment that makes senior workers hard to automate and hollowing them out eliminates the on ramps through which generations first built careers. Widen the on-ramp with apprenticeships and early-career hiring incentives in exposed fields. (We develop this in our response to Question 42.)
- **Steer deployment toward augmentation.** Because firms make the augment-or-displace choice, incentives can tilt it, starting by ending the subsidies in the tax code that make replacing a worker with technology cheaper than keeping one.³²
- **Require disclosure.** The debate has run on anecdote because we cannot see inside firms. Make the measurement of which tasks AI substitutes, augments, and creates permanent and sharper before displacement shows up, not after, so policymakers are working from data rather than press releases.
- **Prepare for wholesale displacement, not just task-trimming.** Hope for augmentation, but build the capacity to respond at scale before it is needed: an unemployment-insurance system that can absorb a fast, concentrated shock; transition support not contingent on a plant-closing-style triggering event; and serious consideration of portable benefits and

wage insurance that share productivity gains rather than letting them pool entirely with capital.

- **Modernize the safety net for task displacement that doesn't rely on mass layoffs.** Task displacement is quieter than whole-job loss and slips through a safety net built for downturns, as nearly 75% of eligible Americans do not even apply for unemployment benefits, and those who do typically receive only 26 weeks of coverage.³³ The fix is to support workers through the transition through portable benefits that follow the worker between jobs, wage insurance that cushions the pay cut when someone is pushed into a lower-paying role, and reskilling paired with real job placement.
- **Take labor-market power seriously.** A worker with options can insist on the augmentation path, while a worker in a concentrated labor market cannot. Scrutiny of employer concentration and non-compete agreements helps determine who captures AI's upside.

Underlying all six is the principle from our agenda: AI-driven productivity gains should be broadly shared, not artificially inflated by avoiding the costs imposed on others. Congress should examine mechanisms to redirect an appropriate share of those gains toward the workers and communities bearing the transition's costs.

Question #42: Protecting entry-level and new-graduate opportunity.

AI is already foreclosing the bottom rungs of the career ladder, as Better Markets documents in its July 2026 report *The State of the Economy for Young Americans*. Entry-level work is the work most susceptible to automation by AI, so the harm concentrates on the young. We're seeing this in the data, as early-career workers (ages 22–25) in AI-exposed occupations have seen a 16% relative decline in employment, and young software developers a nearly 20% drop from their late-2022 peak, even as experienced workers held steady.³⁴ Hollowing out these roles eliminates the entry points through which generations of Americans first built careers. The resulting insecurity is squarely within Better Markets' concern, as it compounds downstream into the low retirement-savings rates, delayed household formation, and speculative "financial nihilism" (crypto, sports betting, prediction markets) that the report documents among young Americans.³⁵

This risk compounds along racial lines. Young workers of color already face higher unemployment and thinner family financial cushions to fall back on during a difficult entry into the labor market. If AI erodes the entry-level roles that have served as the on ramp to the middle class, the workers with the least margin for a delayed or foreclosed start are disproportionately likely to be young people of color. The result would be to widen, at the very first rung, the wealth and opportunity gaps this Committee has long sought to close.

We therefore urge the Committee to treat early-career displacement as a distinct and urgent priority and to respond directly:

- **Widen the on ramp.** Support apprenticeships, paid internships, and early-career hiring incentives in the fields most exposed to task automation, so that the first rung is rebuilt rather than left to wither.

- **Measure the right number.** Require disclosure and measurement that tracks entry-level hiring, not just layoffs, because attrition-plus-hiring-freeze displacement (e.g., a team that keeps two senior editors and simply stops backfilling the ten junior roles below them) is invisible in the monthly jobs report until the category is already hollow.
- **Steer deployment toward augmenting junior workers,** not replacing them, by removing the tax and cost incentives that make automating the entry rung the cheaper path (see our response to Question 41).
- **Investigate directly whether firms capturing AI productivity gains are quietly eliminating the entry-level positions** through which prior generations built careers and whether the resulting insecurity is feeding the retirement-savings and speculative-trading harms Better Markets has documented among young Americans.

AI Data Centers

Question #43: Effects of data centers on grids, water, electricity consumption, and communities.

The costs of the AI buildout are being shifted onto the public, which is the defining concern of the fourth pillar of our agenda.³⁶ Data centers require staggering amounts of electricity: Bloom Energy projects U.S. data-center energy demand will nearly double between 2025 and 2028, and a Bloomberg analysis found wholesale electricity prices near data centers surged as much as 267% over five years.³⁷ In 2025, electric and natural-gas bills became the two largest drivers of inflation, with utility rate-hike requests more than doubling to a record \$31 billion.³⁸ Critically, these costs are regressive, as they fall hardest on lower-income and working-class households who are least likely to benefit from the AI services the infrastructure supports. A Consumer Reports survey found 78% of Americans are concerned data centers will raise their energy bills, and community opposition blocked or delayed twenty data-center projects worth a combined \$98 billion between March and June 2025.³⁹

The environmental and public-health burdens of the buildout—including the strain on local water supplies, emissions from backup generation, and the noise and land-use impacts of large facilities—also tend to concentrate in lower-income communities and communities of color, which have historically borne a disproportionate share of industrial externalities. A framework that requires those capturing AI's gains to bear its costs should focus specifically on whether the burdens are being placed on the communities least positioned to refuse them and least likely to share in the benefits.

Recommendations. The core policy problem is a cost-allocation failure: data centers impose large, concentrated demands on the grid, but under conventional rate design those costs are socialized across all ratepayers, so ordinary households and small businesses subsidize private AI infrastructure. Several states have already begun to correct this, and their approaches offer a

template for federal attention. Virginia's utility regulator approved a dedicated electricity rate class for data centers in November 2025, requiring them to pay for at least 85% of their contracted demand so that other ratepayers are insulated from the cost of capacity built to serve them; Ohio and Texas have enacted comparable measures with legal force.⁴⁰ Building on that experience, we recommend that Congress and the relevant regulators:

- **Make data centers pay for the capacity they require.** Support dedicated rate classes and cost-allocation rules—modeled on Virginia's contracted-demand standard—that place the cost of grid buildout on the large-load customers driving it, rather than on the general ratepayer base. Where interstate wholesale markets are involved, the Federal Energy Regulatory Commission and the regional transmission organizations should ensure these costs are assigned to the loads that cause them.
- **Require transparency as a precondition.** Data-center operators should be required to disclose projected electricity and water consumption and the associated grid and infrastructure costs, so that regulators, ratepayers, and communities can see who will bear them before a project is approved.
- **Protect ratepayers from bearing stranded-asset and speculative-buildout risk.** Because many hyperscalers are not yet profitable and demand projections are uncertain, ratepayers should not be left holding the cost of infrastructure built for demand that may not materialize; long-term take-or-pay commitments and ratepayer safeguards should be structured accordingly. The nonbinding "Ratepayer Protection Pledge" secured from several major technology companies at the urging of the Administration is no substitute for enforceable cost-allocation rules.⁴¹
- **Recognize the political salience and act rather than defer.** With lawmakers in more than 25 states already legislating on data-center cost allocation, moratoriums, and water reporting, federal inaction cedes the field and risks an inconsistent patchwork; a floor of federal ratepayer protections would complement, not displace, state action (see our response to Question 46).

Question #44: Structure of tax breaks and financing mechanisms for data centers.

Better Markets has not yet published a detailed technical analysis of the specific tax and financing structures underlying individual data-center deals, so we defer to commenters with transaction-level expertise on the mechanics. We note only the policy consequences our work does address: to the extent state and local tax breaks and public financing subsidize private AI infrastructure, they represent a further channel through which ordinary ratepayers and taxpayers are subsidizing private AI profits.

Question #45: Ways taxpayers are subsidizing data centers and AI development.

The subsidy runs through multiple channels: ratepayer-funded grid buildouts to serve data-center demand, state and local tax incentives, and the environmental and public-health externalities



communities absorb. The unifying question from our agenda is whether ordinary ratepayers are subsidizing private AI profits and, more broadly, whether those benefiting from technological transformation artificially inflate their gains by shifting costs onto families, communities, ratepayers, and taxpayers. Better Markets recommends that these subsidies be made visible and measurable as a precondition to any judgment about whether they serve the public interest.

Question #46: Federal safeguards in line with state and local moratoriums.

Communities are not waiting: lawmakers in more than 25 states introduced data-center bills in 2026, and some jurisdictions have declared moratoriums.⁴² Better Markets supports federal attention to ensuring that the costs of the buildout are borne by those capturing its gains. Whatever the specific mechanism, the governing principle should be that AI infrastructure is not permitted to inflate private profits by externalizing its costs onto the public.

Question #47: Community benefits agreements.

Better Markets supports community benefits agreements as one promising tool, while cautioning that they are not a substitute for adequate regulation. When companies build massive data centers in communities, they incur real obligations to those communities, and local governments are too often not equipped to negotiate fair terms. We recommend that any community-benefits framework address the concrete, measurable costs our research identifies (e.g., electricity-rate impacts, water use, and environmental and public-health burdens) and that federal policy help correct the negotiating-power imbalance between global technology firms and the local governments across the table from them.

Financial Stability

Question #49: Steps FSOC, OFR, and banking regulators should take using existing authorities.

AI's financial-stability risks such as herding, correlated failure, and systemic contagion, are increasingly recognized, as FSOC first identified AI as a vulnerability in the financial system in its 2023 annual report.⁴³ The mechanism is well understood: when multiple institutions rely on similar AI-driven strategies trained on the same data, they produce correlated behavior that amplifies market shocks rather than absorbing them. Agentic AI compounds this, exhibiting emergent behaviors—including collusion and manipulation—that traditional oversight is not designed to detect.⁴⁴

Regulators can act now, within existing authorities, to (a) monitor the concentration of common AI models and data dependencies across institutions as a systemic-risk vector (the “monoculture of models” concern), (b) incorporate AI-driven correlated behavior into stress testing and scenario analysis, and (c) direct the OFR to build the data and analytical capacity to detect emergent, cross-institutional AI risks before they cascade.

Question #50: Congressional steps to improve authorities or modify laws for AI financial-stability risk.

Where existing authorities prove insufficient, Congress should act to (a) ensure FSOC has clear authority and the will to designate systemically important nonbank AI dependencies (Better Markets has long noted the anomaly that not a single nonbank financial institution is currently designated systemically important despite the government's repeated interventions to prop up markets)⁴⁵ and (b) close the third-party-vendor oversight gaps described in our response to Question 20, so that critical AI service providers to the financial system do not sit beyond the reach of supervision.

Question #51: Regulators' own use of AI to monitor the financial system.

Better Markets supports regulators using AI to strengthen supervision and market monitoring, with two conditions. First, meaningful human control and accountability must be preserved at every consequential stage. An AI tool may surface a pattern, but human judgment must weigh it. Second, regulators' AI capacity must not become a channel for industry capture. The models, data, and personnel must serve the public-interest mission, not the perspectives of the firms being supervised. Used well, regulatory AI (i.e., SupTech) could help level a playing field between industry's mighty shadow and the public interest.



Conclusion

The questions in this RFI are the right questions, and the Committee deserves credit for asking them while the choices are still open. The consistent thread in Better Markets' answers is that AI's risks in finance are neither speculative nor distant: discriminatory AI underwriting has already produced settlements and audit findings, AI credit models are already producing denials their own deployers cannot explain, algorithmic trading already dominates markets and has already produced accountability-free cascades, and the costs of the AI buildout are already landing on the electricity bills of families who see little of the benefit.

Much of what is needed does not require Congress to write on a blank slate. It requires that existing anti-discrimination, consumer-protection, securities, and safety-and-soundness laws be applied to AI and enforced. It also requires that the agencies charged with that enforcement (above all the CFPB) be preserved rather than dismantled. Finally, it also requires that the model-risk framework cover the models that actually pose the risk, that liability be made clear before these systems entrench, and that the winners of the AI transition not be permitted to inflate their gains by shifting the costs onto everyone else. And because AI learns and scales the disparities embedded in historical data, applying and enforcing the nation's fair-lending and fair-housing laws in the AI era is not a peripheral concern, it is central to ensuring that this technological transition narrows rather than widens the racial wealth gap.

Change is certain. Progress is not. Better Markets stands ready to work with the Committee to ensure that, this time, the public interest has a seat at the table before the rules harden.

Respectfully submitted,

Better Markets, Inc.



Better Banks | Better Businesses
Better Jobs | Better Economic Growth
Better Lives | Better Communities

Better Markets is a public interest 501(c)(3) non-profit based in Washington, DC that advocates for greater transparency, accountability, and oversight in the domestic and global capital and commodity markets, to protect the American Dream of homes, jobs, savings, education, a secure retirement, and a rising standard of living.

Better Markets fights for the economic security, opportunity, and prosperity of the American people by working to enact financial reform, to prevent another financial crash and the diversion of trillions of taxpayer dollars to bailing out the financial system.

By being a counterweight to Wall Street's biggest financial firms through the policymaking and rulemaking process, Better Markets is supporting pragmatic rules and a strong banking and financial system that enables stability, growth, and broad-based prosperity. Better Markets also fights to refocus finance on the real economy, empower the buy-side and protect investors and consumers.

For press inquiries, please contact us at press@bettermarkets.org or (202) 618-6430.



[SUBSCRIBE](#) to Our Monthly Newsletter

FOLLOW US ON SOCIAL



2000 Pennsylvania Avenue NW | Suite 4008 | Washington, DC 20006 | 202-618-6464 | www.bettermarkets.org
© 2026 Better Markets, Inc. All Rights reserved.

-
- ¹ Better Markets, “[AI, the Economy, and You: A People-Centered AI Agenda](#)” (June 2026); Better Markets, “[Who We Are](#)”.
- ² Better Markets Substack, “[What We Know about the AI and Jobs Debate, and Who’s in Its Crosshairs](#)” (July 2026)
- ³ Better Markets, “[The CFPB Is an Irreplaceable Consumer Protection Agency](#)” (March 2025), “[The Trump Administration’s Attack on Consumers in Five Facts](#)” (July 2025), and “[The Demise of Consumer Financial Protection Regulations under Trump’s CFPB](#)” (July 2025).
- ⁴ Federal Reserve, “[SR 26-2: Revised Guidance on Model Risk Management](#)” (April 2026).
- ⁵ Consumer Financial Protection Bureau, “[CFPB Issues Guidance on Credit Denials by Lenders Using Artificial Intelligence](#)” (September 2023); [15 U.S.C. §§ 1691–1691f \(ECOA\)](#).
- ⁶ For example, see Massachusetts Office of the Attorney General, “[AG Campbell Announces \\$2.5 Million Settlement With Student Loan Lender For Unlawful Practices Through AI Use, Other Consumer Protection Violations](#)” (July 2025) and Relman Colfax, “[Fair Lending Monitorship of Upstart Network’s Lending Model \(Final Report\)](#)” (March 2024).
- ⁷ AICerts News, “[Automated Trading Risk Exposed in Crypto Flash Crash](#)” (October 2025).
- ⁸ CFPB, “[CFPB Issues Guidance on Credit Denials by Lenders Using Artificial Intelligence](#)” (September 2023)
- ⁹ CFPB, “Supervisory Highlights: Advanced Technologies Special Edition,” (January 2025). While the CFPB no longer hosts this document on its website, a copy can be found [here](#).
- ¹⁰ Massachusetts Office of the Attorney General, “[AG Campbell Announces \\$2.5 Million Settlement With Student Loan Lender For Unlawful Practices Through AI Use, Other Consumer Protection Violations](#)” (July 2025).
- ¹¹ Relman Colfax, “[Fair Lending Monitorship of Upstart Network’s Lending Model \(Final Report\)](#)” (March 2024).
- ¹² Executive Office of the President, “[Executive Order 14281: Restoring Equality of Opportunity and Meritocracy](#)” (April 2025).
- ¹³ Federal Reserve, “[SR 26-2: Revised Guidance on Model Risk Management](#)” (April 2026).
- ¹⁴ Federal Reserve and the Office of the Comptroller of the Currency, “[SR 11-7: Guidance on Model Risk Management](#)” (April 2011).
- ¹⁵ Securities and Exchange Commission, “[SEC Charges Two Investment Advisers with Making False and Misleading Statements About Their Use of Artificial Intelligence](#)” (March 2024).
- ¹⁶ Better Markets, “[SEC’s Proposal to Reduce the Frequency of Public Company Disclosures Threatens Investor Protection](#)” (May 2026).
- ¹⁷ U.S. House Committee on Financial Services Democrats, “[Ranking Member Waters, Chair McHenry, Representatives Hill and Lynch Release Bipartisan AI Working Group Staff Report](#)” (July 2024).
- ¹⁸ Better Markets, “[The Trump Administration’s Attack on Consumers in Five Facts](#)” (2025); Better Markets, “[The CFPB Is an Irreplaceable Consumer Protection Agency](#)” (Aug. 25, 2025).
- ¹⁹ For example, see Better Markets, “[Community Banks: Vital to Main Street Families, Small Businesses, and Local Economies](#)” (April 2025), “[Strengthening Community Banks Creates an Economy That Works for All Americans](#)” (September 2025), “[Community Banks Need More Than the Crumbs from Big Banks’ Cake to Thrive](#)” (May 2026), and “[Don’t Let Private Equity Destroy Community Banks](#)” (April 2026).
- ²⁰ For example, see CBS News, “[UnitedHealth uses faulty AI to deny elderly patients medically necessary coverage](#)” (November 2023).
- ²¹ AICerts News, “[Automated Trading Risk Exposed in Crypto Flash Crash](#)” (October 2025).
- ²² [Moffatt v. Air Canada, British Columbia Civil Resolution Tribunal](#) (February 2024).
- ²³ [15 U.S.C. §§ 1691–1691f \(ECOA\)](#).
- ²⁴ Better Markets, “[AI, the Economy, and You: A People-Centered AI Agenda](#)” (June 2026).
- ²⁵ Better Markets, “[SEC’s Proposal to Reduce the Frequency of Public Company Disclosures Threatens Investor Protection](#)” (May 2026).
- ²⁶ AICerts News, “[Automated Trading Risk Exposed in Crypto Flash Crash](#)” (October 2025).
- ²⁷ Oxford University, “[The Agentic Regulator: Risks for AI in Finance and a Proposed Agent-based Framework for Governance](#)” (December 2025).
- ²⁸ SEC, “[SEC Charges Two Investment Advisers with Making False and Misleading Statements About Their Use of Artificial Intelligence](#)” (March 2024)
- ²⁹ California Law Review, “[A Home For Digital Equity: Algorithmic Redlining and Property Technology](#)” (October 2023)
- ³⁰ Executive Order: “[Restoring Equality of Opportunity and Meritocracy](#)” (April 2025); “[Mass. AG, Earnest Settlement](#)” (July 2025); “[Relman Colfax, Upstart Monitorship Final Report](#)” (March 2024).
- ³¹ U.S. Census Bureau, “[The Microstructure of AI Diffusion: Evidence from Firms, Business Functions, and Worker](#)

[Tasks](#)” (April 2026).

³² See Better Markets Substack, “[The Tax Code Has Its Thumb on the Scale](#)” (July 2026).

³³ Fortune, “[AI layoffs are coming. The problem may be compounded because nearly 75% of people don’t apply for unemployment benefits](#)” (March 2026).

³⁴ Brynjolfsson, Chandar & Chen, “[Canaries in the Coal Mine? Six Facts about the Recent Employment Effects of Artificial Intelligence](#)” (November 2025).

³⁵ See Better Markets, “[The State of the Economy for Young Americans](#)” (July 2026).

³⁶ Better Markets, “[AI, the Economy, and You: A People-Centered AI Agenda](#)” (June 2026).

³⁷ Bloom Energy, “[2026 Data Center Power Report](#)” (January 2026); Bloomberg News, “[AI Data Centers Are Sending Power Bills Soaring](#)” (September 2025).

³⁸ Fortune, “[Your utility bills keep going up. Here’s everyone you can blame—AI data centers included](#)” (March 2026).

³⁹ Consumer Reports, “[American Experiences Survey](#)” (November 2025); MultiState, “[State Data Center Laws Challenge Federal AI Infrastructure Push](#),” (April 2026).

⁴⁰ American Action Forum, “[Virginia’s New Data Center Electricity Rate Class](#)” (January 2026); Statehouse News Bureau, “[Ohio House bill would extend data center tariff to the rest of state](#)” (March 2026); Office of the Texas Governor, “[PUC and ERCOT Actions Protecting Residential Ratepayers](#)” (July 2026).

⁴¹ White House, “[Ratepayer Protection Pledge](#)” (March 2026).

⁴² MultiState, “[State Data Center Laws Challenge Federal AI Infrastructure Push](#),” (April 2026).

⁴³ U.S. Department of the Treasury, “[FSOC 2023 Annual Report](#)” (December 2023).

⁴⁴ Oxford University, “[The Agentic Regulator: Risks for AI in Finance and a Proposed Agent-based Framework for Governance](#)” (December 2025).

⁴⁵ Better Markets, “[Community Banks: Vital to Main Street Families, Small Businesses, and Local Economies](#)” (April 2025).