

NBER WORKING PAPER SERIES

THE COASEAN SINGULARITY? DEMAND, SUPPLY, AND MARKET DESIGN
WITH AI AGENTS

Peyman Shahidi
Gili Rusak
Benjamin S. Manning
Andrey Fradkin
John J. Horton

Working Paper 34468
<http://www.nber.org/papers/w34468>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
November 2025

All authors contributed equally. Fradkin and Horton use seniority-based ordering. Fradkin is currently employed by Amazon. The views expressed in this chapter are solely those of the authors and do not reflect the views of Amazon Inc. or the National Bureau of Economic Research.

At least one co-author has disclosed additional relationships of potential relevance for this research. Further information is available online at <http://www.nber.org/papers/w34468>

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2025 by Peyman Shahidi, Gili Rusak, Benjamin S. Manning, Andrey Fradkin, and John J. Horton. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

The Coasean Singularity? Demand, Supply, and Market Design with AI Agents
Peyman Shahidi, Gili Rusak, Benjamin S. Manning, Andrey Fradkin, and John J. Horton
NBER Working Paper No. 34468
November 2025
JEL No. D47, D83, J44, K20, L15, L86, O33

ABSTRACT

AI agents—autonomous systems that perceive, reason, and act on behalf of human principals—are poised to transform digital markets by dramatically reducing transaction costs. This chapter evaluates the economic implications of this transition, adopting a consumer-oriented view of agents as market participants that can search, negotiate, and transact directly. From the demand side, agent adoption reflects derived demand: users trade off decision quality against effort reduction, with outcomes mediated by agent capability and task context. On the supply side, firms will design, integrate, and monetize agents, with outcomes hinging on whether agents operate within or across platforms. At the market level, agents create efficiency gains from lower search, communication, and contracting costs, but also introduce frictions such as congestion and price obfuscation. By lowering the costs of preference elicitation, contract enforcement, and identity verification, agents expand the feasible set of market designs but also raise novel regulatory challenges. While the net welfare effects remain an empirical question, the rapid onset of AI-mediated transactions presents a unique opportunity for economic research to inform real-world policy and market design.

Peyman Shahidi
Massachusetts Institute of Technology
Sloan School of Management
shahidi.peyman96@gmail.com

Gili Rusak
Harvard University
gilirusak@g.harvard.edu

Benjamin S. Manning
Massachusetts Institute of Technology
bmanning@mit.edu

Andrey Fradkin
Boston University
fradkin@bu.edu

John J. Horton
Massachusetts Institute of Technology
MIT Sloan School of Management
and NBER
john.joseph.horton@gmail.com

1 Introduction

AI agents are autonomous software systems that perceive, reason, and act in digital environments to achieve goals on behalf of human principals, with capabilities for tool use, economic transactions, and strategic interaction.¹ This chapter focuses on the transformative implications of these systems as market participants. We envision agents as having the ability to harness computational resources, to communicate with other agents and humans, to receive and send money, and to access and interact with the Internet.^{2,3} A variety of agents are already available (e.g., AutoGPT, Replika, Claude Code), and more capable ones are in the development pipeline.

A prototypical example of an AI agent operating autonomously is *Deep Research* (Citron 2024), introduced by Google’s Gemini team. Unlike traditional software, which retrieves information or processes only pre-supplied data, *Deep Research* takes natural language instructions (prompts) and independently carries out actions to produce a researched report. An economist writing a paper without such an agent might rely on Google Scholar to find articles, refine queries, and synthesize results—where retrieval is automated but reasoning remains with the user. Or they might upload papers into an AI system to generate a summary, which still confines the system to provided inputs.⁴ In both cases, the software does not autonomously define or pursue tasks. By contrast, *Deep Research* can iteratively search the web, evaluate results, and assemble a report without human oversight at each intermediate step. This ability to perceive, reason, and act on natural language instructions is what makes AI agents distinctive.

We offer a practical perspective on AI agents as participants in digital markets, complementing Hadfield and Koh’s (2025) more theoretical treatment in this volume. We take seriously the idea that, rather than hiring a human agent, one could instead rely on an AI agent. This possibility is near, as agents already perform, albeit imperfectly, tasks such as shopping (Dammu et al. 2025), negotiation (Zhu et al. 2025), or search for products on e-commerce sites (Zeff 2025; Allouah et al.

¹We use the term “principal” to mean any stakeholder that deploys AI agents. The principal can be a consumer/firm that hires an assistant agent to act on their behalf in a market activity, or a business/platform that deploys service agents to interact with consumers or with consumers’ representative agents.

²Throughout this chapter, we use “agents” and “AI agents” interchangeably, except where we explicitly refer to human agents.

³For the purposes of this essay, we abstract away from Artificial Super Intelligence (ASI), in which AI is omniscient in every dimension relative to humans.

⁴Currently, the term “Large Language Model” is commonly used to describe AI models with advanced capabilities. However, we avoid using this term throughout the essay because frontier AI architectures and their popular names will likely evolve over time. Instead, we opt for “AI agent” or “AI system.”

2025). It is easy to imagine more advanced agents contacting multiple counterparties to negotiate prices or applying to jobs and advocating for employment on a user’s behalf.

The fundamental economic promise of AI agents lies in their ability to dramatically reduce transaction costs—the expenses associated with using markets to coordinate economic activity. This reduction occurs not only through direct task execution but also via advisory services. In the direct approach, agents perform tasks entirely on behalf of users, from price comparison shopping to contract negotiation to job interviews. In the advisory approach, agents assist users in making better market decisions, such as helping job applicants individually optimize their resumes for each job submission. Recent evidence suggests both approaches can be quite effective. For the former, at least one study has shown that outcomes can be better for job seekers interviewed by AI as opposed to human employers (Jabarian and Henkel 2025). For the latter, algorithmic assistance on job application essays can increase hiring rates significantly (Wiles et al. 2025).

It is important to recognize that demand for AI agents represents derived demand rather than direct utility. Individuals do not generally derive satisfaction from watching an agent compile price lists for gas grills; they employ the agent purely to achieve some desired market outcome following their own decision-making process. As a result, humans will deploy AI agents in two primary scenarios: first, to optimize decisions that they would otherwise make sub-optimally due to information constraints or cognitive limitations; and second, to make decisions of similar or potentially even lower quality, but at dramatically reduced cost and effort.

The broad adoption of AI agents will cause transformative downstream effects on the economy. Although the exact nature of these changes remains uncertain, the forces at play are familiar to economists: supply and demand of AI agents will continue to shape market organization, while technological change alters the relative costs of different activities. To see how these forces might operate, Coase’s (1937) insight that transaction costs play a central role in shaping organizations is particularly useful. One could argue that much of how we structure our economy and firms can be explained by transaction costs, often costs of human labor. The activities that comprise transaction costs—learning prices, negotiating terms, writing contracts, and monitoring compliance—are precisely the types of tasks that AI agents can potentially perform at very low marginal cost. Once agents can indeed execute these functions effectively and cheaply, we will see significant shifts in the traditional make-or-buy boundaries that define firm organization and market structure.

Transformations in market structure will not only come from existing markets adapting to agent capabilities but also from the emergence of entirely new, agent-first designs. In markets that adapt, agents will initially augment humans by assisting with specific tasks, then, as their capabilities grow, will begin to substitute for tasks entirely, shifting human effort toward judgment, oversight, and relationship-focused activities. Over time, this progression will lead to reorganization, as tasks are removed or redefined and workflows are increasingly structured around agent capabilities. In contrast, novel, agent-first markets will start from that endpoint. Rather than evolving through incremental augmentation and substitution, they will be designed from the ground up with processes, interactions, and roles structured entirely around agent capabilities.

AI agents expand the economic market design frontier. By collapsing the costs of eliciting preferences, enforcing commitments, and verifying identity, agents make mechanisms that were once only theoretically attractive now feasible at scale. This new set of feasible designs has the potential to improve consumer welfare and matching. That said, AI agents are not guaranteed to make markets more efficient or to improve other social objectives. Even if it is individually rational for consumers and firms to adopt agents, the equilibrium outcomes may be suboptimal. This is particularly relevant in environments with externalities across agents, imperfect alignment, and asymmetric information. These factors point to an exciting market design agenda for economists: translating theoretical insights into practical mechanisms that can guide the transition to agent-first environments and capture the full benefits of AI.

The rest of this chapter proceeds as follows. Section 2 analyzes demand for AI agents. Section 3 considers agent design considerations. Section 4 discusses agent supply. Section 5 addresses equilibrium effects of agent adoption under the status quo. Section 6 focuses on transformative changes that the proliferation of AI agents will have on markets. Section 7 overviews some regulation issues concerning agents. Section 8 concludes.

2 Demand for AI Agents

Humans will demand AI agents for the same reasons they demand human agents: they believe it is rational to have some aspect of a decision or market transaction handled through an intermediary. This might occur because the agent’s time is less expensive than the principal’s time, because the

agent is legitimately better at the task than the principal, or because the principal has reasons for obscuring their identity in the transaction.

These motivations will increasingly drive the adoption of AI agents along two avenues. First, agents will substitute for human intermediaries where one would otherwise do the work personally or hire a human agent. The canonical example is product search. AI agents convert the costly, time-consuming parts of intermediation—search, screening, quoting, negotiation, scheduling—into low-cost compute and API calls. An agent can solicit and compare many quotes in parallel, then book and monitor follow-through at far lower marginal cost than a human. Even when skilled execution remains human, the intermediation premium falls, creating demand for AI agents.

Second, and perhaps more consequential, agents will enable undertakings that would not have been attempted at all. By lowering the cost of exploration and execution, they expand the feasible set of options and reduce the threshold of “worth doing.” For physical jobs (e.g., home installation, fixing a sink), the agent can conduct diagnostic triage, source parts, select vendors, and schedule service. For software, it can generate and iterate on a bespoke script. Compared to humans, agents can persist through repeated failures at much lower marginal cost and continue monitoring tasks over longer windows to secure better outcomes.

AI agents are likely to first gain traction in markets where human agency is already common. Key characteristics of these markets include high-stakes transactions, large pools of potential counterparties, substantial effort required to evaluate options, information asymmetries that could be resolved through effort or experience, and experience asymmetries due to differences in transaction frequency. Examples of such markets include job search, real estate, and certain investment contexts where counterparties matter beyond financial terms alone. A useful heuristic suggests AI agents will appear first in markets that already employ human agents or where large digital platforms were created to overcome matching frictions—such as LinkedIn, Upwork, Zillow, and Airbnb. Table 1 summarizes these characteristics and associated markets, highlights the existing human or platform-based solutions that address them today, and shows how AI agents can improve on those approaches. As adoption spreads from these initial footholds, the dynamics of trust, evaluation, and inter-agent coordination will shape broader economic integration.

Although the reasons for the demand for AI agency are clear, how this will translate into a willingness to pay is less clear. Consumers will want the same core attributes that they seek in

Table 1: Where AI Agents First Gain Traction

Market Characteristic	Example Markets	Existing Solutions	How AI Agents Help
High-stakes Transactions	Real estate, Job search, Investment decisions	Human agents (realtors, headhunters, financial advisors)	AI agents can analyze vast amounts of data and documentation without fatigue, providing thorough due diligence at near-zero marginal cost.
Vast Counterparty Space	Dating, Freelance hiring, Rental markets	Digital platforms (Tinder, Upwork, Airbnb)	AI agents can evaluate thousands of options simultaneously, with no opportunity cost to their “time”—they can search exhaustively where humans must sample.
High Evaluation Effort	Startup funding, College admissions, B2B procurement	Specialized consultants, Matching services	AI agents can read every review, analyze every metric, and compare all attributes across options without the time constraints that force humans to use heuristics.
Information Asymmetries	Used car markets, Insurance shopping, Legal services	Brokers, Comparison sites, Expert intermediaries	AI agents can continuously monitor multiple information sources, cross-reference data, and identify discrepancies that would take humans hours to uncover.
Experience Asymmetries	Home buying (once per decade vs. daily), Wedding planning, Estate planning	Professional agents who transact frequently	AI agents can leverage collective experience from millions of transactions, effectively giving every user the negotiating power of a frequent transactor.

human agents: capability sufficient to act on their preferences successfully, knowledge of their preferences, and alignment sufficient to act on their preferences to their benefit. In essence, they will want capable, knowledgeable, and faithful agents. Demand for reliable information about agent performance will be substantial, with benchmarks playing important roles alongside word-of-mouth recommendations, even as consumers rightly worry about benchmark gaming. Differentiation will

likely emerge, with some agents becoming known as particularly effective for specific applications. Yet ensuring that agents consistently deliver on these characteristics poses substantial theoretical and practical challenges—ones that future research will need to address.

3 Designing AI Agents

In the previous section, we explored why humans will demand AI agents. But of course, this product/service does not yet exist, at least in the form being imagined. In this section, we consider the design and development of agents. While we cannot hope to describe the precise design, we can speculate on the key challenges and the focus of R&D efforts. The use of AI agents will be driven by their design, which has both an engineering and an economic component. On the engineering side, there are practical challenges in having capable agents that can interact with the digital world in a reliable manner. We set aside the engineering challenges for the purposes of this discussion, though they are substantial in their own right (Kalai et al. 2025). Instead, we focus on the economic component, namely what actions should a capable agent take.

Clearly, an AI agent must know the principal’s preferences to act on their behalf. The core design problem is thus two-fold: both eliciting those preferences and ensuring the agent honors them—together these are what computer scientists and philosophers refer to as the alignment problem (Bostrom 2014, Christian 2020). Note that this encompasses the economic principal-agent problem, where preferences are already known and the challenge is designing contracts or incentives to prevent self-interested agents from shirking their duties.

Preference elicitation, in particular, offers new challenges even though the practice itself is not new to digital markets. Conventional recommendation systems already translate billions of online consumer choices into predicted preferences using complex machine learning algorithms. While these traditional systems are powerful, their flexibility remains fundamentally limited. They operate with fixed input and output dimensionalities and are trained for specific contexts. For example, Netflix’s current television recommendation algorithm cannot recommend consumer or financial goods—even if streaming choices contain latent information about such products.

AI agents, by contrast, operate in natural language and other open-ended mediums rather than within fixed input structures. This makes them far more flexible: any statement or request

can, in principle, serve as input. Flexibility, however, can lead to unexpected outcomes because it is not possible for a principal to fully specify all edge cases an agent might encounter. As such, failures to recover a principal’s “full” preferences can still arise for two distinct reasons: (1) principal’s articulation limits: the principal cannot fully or consistently specify their preferences (e.g., providing a rank-ordering or pairwise comparisons of the entire Netflix catalog is impractical, even if an individual could generate such an ordering), and (2) agent’s synthesis limits: the agent misinterprets what is stated, including via inaccurate inference or hallucination.

Importantly, the complexity or dimension of human preferences varies dramatically across domains. In some cases, human preferences are relatively straightforward: a home seller wants to maximize price while minimizing time to sale, requiring the agent to understand only the speed-price trade-off. By contrast, a home buyer’s preferences are far more complex. A buyer may care about dozens of factors simultaneously—location, commute time, school quality, neighborhood safety, size and layout of the house, style, age of the property, price, and future resale value, among others. It is correspondingly easier for most people to decide whom to sell their house to rather than which house to buy; the input to the former is of much lower dimensionality than the latter. When preferences are high-dimensional, even small errors in reporting can propagate (Liang 2025), making the underlying alignment task inherently more difficult. Thus we expect that the dimensionality of preferences within a context is a key predictor of agent adoption and thus demand.

Beyond preference elicitation, a central design challenge concerns meta-rationality: agents must learn not only how to act, but also when to act autonomously and when to defer to their principals. Determining these boundaries is essential for preserving trust, ensuring that delegation improves welfare rather than creating new risks of overreach. Closely related is the need for agents to be both rational and robust. Rationality here refers to consistency in decision-making under uncertainty, while robustness entails resilience against adversarial manipulation. As agents become more integral to market interactions, incentives to exploit their decision rules will intensify, raising concerns analogous to adversarial attacks in other domains (e.g., “black-hat” optimization of content for ranking systems).

Taken together, these challenges suggest a broader research program at the intersection of economics and computer science: how to design agents that are rational, aligned, resistant to

manipulation, and capable of striking an appropriate balance between autonomy and deference.

4 Supply of AI Agents

We consider the supply of AI agents from two complementary perspectives. First is the production side, where firms use foundation models to source agents and then choose to price them.⁵ Second is the consumer side, which concerns how users experience agents in market transactions.

To understand the production side, we begin with the technological innovation underlying nearly all agents: foundation models. These are the workhorses behind AI agents. Already two types of agent providers exist in terms of model use: those who build their own foundation models (e.g., Anthropic, OpenAI), and those who use others' foundation models and customize them for particular use cases (e.g., Decagon, Harvey, Sierra). The economics of these differs sharply, since training foundation models incurs high fixed costs, while operating agents incurs mostly variable costs.

The equilibrium structure of the agent industry will be shaped not only by development costs but also by strategic platform behavior. Decisions such as whether to interoperate with or exclude rival agents will affect the scope of competition, concentration, and ultimately how surplus is divided. It remains an open question whether successful agent providers will end up building their own foundation models, either to exploit returns to scale or to capture scope advantages from vertical integration. Control over these models and the terms of access directly shapes which firms can supply agents, how they differentiate, and how competition unfolds. If vertical integration proves essential, concentration among agent providers is likely to mirror the concentration already present in the foundation model industry (Fradkin 2025). Even when access is open, providers gain advantage by refining model use and developing complementary agent capabilities using proprietary datasets, which offers another lever for differentiation.

Another equally important dimension of supply is agent pricing. Human agent services are often priced as a percentage of the transaction value. The sums involved can be high in situations where the human agent is paid to engage in an adversarial situation against another human agent, and where one side wins and another loses. These high sums in adversarial settings are due to two

⁵A foundation model is a large-scale AI model trained on broad, diverse datasets that serves as a base for multiple downstream applications.

types of market forces: that the “best” human agents are scarce and that high-powered incentives result in extraordinary effort.

AI agents, however, face a very different cost and supply profile, which changes the pricing dynamics. Because software can be copied at negligible cost, the supply of AI agents will be virtually unconstrained. Unlike human agents, they also do not derive utility from pecuniary compensation. These factors point to diminished pricing power for AI relative to human agents. However, an opposing force pushing toward higher agent prices arises if the quality of an agent’s performance depends on the compute allocated to it. In that case, prices could scale with the amount of compute required—potentially rising in proportion to the stakes of the transaction.

Even if performance improves with compute, the returns to agent quality are unlikely to increase indefinitely. As marginal gains diminish, the market dynamics begin to resemble those of other digital services such as search and social media. In this scenario, the agent may be offered for free and supported by advertising, or bundled with complementary goods (e.g., phones) and services (e.g., delivery). Providers may also experiment with more sophisticated models, including tiered pricing—where a limited free version serves price-sensitive users while a premium plan with full functionality cross-subsidizes it—or two-part tariffs, in which users pay an upfront subscription plus per-prompt or per-token fees.

While industry dynamics determine who produces agents and how they are priced, they ultimately manifest for consumers through the types of agents they can access and the terms of use. In particular, we envision consumers experiencing different types of agents along two key dimensions: ownership, which concerns who has custody of the agent in a market transaction, and specialization, which reflects the breadth or narrowness of the agent’s capabilities in a specific context. These two axes together yield four archetypes of AI agent supply, summarized in Table 2.

In terms of ownership, consumers will face two types of agents: “bring-your-own” and “bowling-shoe.” Presumably, few consumers will literally create their own models so they will rely on firms to supply them.⁶ Consider, for example, an Anthropic-powered agent being used on Walmart and Amazon via public APIs. This is a bring-your-own agent: it carries the user’s instructions and data across sites, and neither platform sees or dictates its programming beyond what it explicitly

⁶By “their own” we mean an agent created by some marketplace or platform that is itself a participant in an exchange.

shares. By contrast, a bowling-shoe agent is provided by the platform hosting the transaction. These agents enjoy deep integration, privileged signals, and low setup costs, but are less portable, may steer outcomes toward the platform’s interests, and can contribute to platform lock-in.

Orthogonal to ownership is specialization. We posit that there are likely to be two types of AI agents along this dimension: horizontal and vertical. Horizontal agents are generalists that can span many tasks and platforms, leveraging a single memory/preferences layer across markets and arbitraging options when beneficial, but they may lack domain-specific tooling or compliance features. Vertical agents, on the other hand, specialize in a narrow domain (e.g., tax filing, job search, travel) or even a specific platform workflow, trading breadth for depth (i.e., stronger performance, richer integrations, and tailored guardrails) at the cost of portability and reuse.

Table 2: Ownership and Specialization Dimensions of AI Agents

Ownership	Specialization	
	Horizontal	Vertical
Bring-your-own Agent	Features: User-controlled agent; not operated by the platform; carries cross-site memory/preferences; uses public APIs/standard interfaces; limited privileged hooks. Pros: Portable across markets; strong user alignment; privacy; can compare/arbitrate across platforms. Cons: Possible throttling/degraded access; weaker first-party data/tools; setup/subscription/compute cost.	Features: User-controlled specialist for a narrow domain (e.g., tax, jobs, travel); interoperates across competing platforms within that domain; third-party (not platform-run). Pros: Higher task performance than BYO-horizontal; retains user alignment; reusable across multiple platforms in-domain. Cons: Still lacks platform-privileged integrations; must track per-platform APIs/policies; limited reuse outside the domain.
Bowling-shoe Agent	Features: Platform-operated generalist embedded in OS/app/site; convenient defaults; first-party telemetry and UI control. Pros: Low user friction; strong latency/reliability; access to proprietary features/tools. Cons: Limited portability; steering/self-preferencing risk; lock-in; weaker inspectability.	Features: Platform-operated specialist tightly integrated with domain tooling, policies, and datasets; optimized end-to-end flows with guardrails/compliance. Pros: Best performance on owning platform; richest domain functionality; full UX control for the platform. Cons: Highest degree of steering/lock-in; least transparent/inspectable; cross-platform substitution discouraged.

The choice between bring-your-own-agent and platform-provided agent models presents trade-

offs from the consumer perspective. Bringing your own agent offers the advantages of perfect (or at least better) alignment with personal preferences and privacy, cross-platform functionality, and access to detailed personal information, allowing users to maintain consistent experiences across different services. However, this approach comes with significant maintenance costs and the risk of being outperformed by more sophisticated platform-specific alternatives. Ultimately, platforms might even restrict direct access to their services, requiring users to go through their preferred agent intermediaries (Rothschild et al. 2025). Conversely, platform-provided agents eliminate technical complexity and may be better if they are trained on platform specific data. But bowling-shoe agents may also sacrifice perfect alignment with individual needs, either due to explicit self-preferencing or due to simply not considering options offered on other platforms.

From a platform provider’s perspective, the choice of AI agent model creates strategic trade-offs around control, costs, and competitive positioning. The bring-your-own-agent model reduces computational and hosting expenses for platforms while minimizing liability and simplifying API maintenance, but it sacrifices usage insights, opportunities for optimizing the user experience, and potentially profitable opportunities to steer consumers to the platform’s preferred options. Furthermore, the bring-your-own-agent model reduces platform lock-in effects.

Platform-provided agents enable companies to maintain control over the user experience, benefit from their own R&D investments, and potentially steer users toward preferred options, but require substantial hosting and computation resources while potentially suppressing consumer demand due to alignment concerns. This fundamental tension between user autonomy and platform control mirrors broader debates in digital markets about the optimal balance between personalization, cost, and market power.

Looking ahead, we can also imagine more complex structures, such as Anthropic’s horizontal agent being integrated with an iPhone, or asking for help from Walmart’s vertical agent, or interacting directly with a seller’s agent while dis-intermediating Walmart altogether. Agent-to-agent interfaces and agent-only storefronts may also proliferate. Interfaces that are made primarily for agents may be useful because of their speed, and the ability of the interface provider to tailor information to the agent. From a consumer’s perspective, however, the agentic interaction becomes less inspectable—highlighting transparency and trust as the critical challenges that will shape adoption.

5 Equilibrium Effects of AI Agency Under the Status Quo

To understand how AI agents will reshape markets, we can examine their effects within existing market structures before considering how market design itself might evolve. Consider a prototypical e-commerce market where agents enable users to consider all available options with complete information access. In such environments, sophisticated agents would likely prove resistant to nudges and advertising that lacks informational or signaling value, fundamentally altering competitive dynamics. These mechanisms highlight substantial efficiency gains but also introduce new risks and distributional uncertainties, leaving overall welfare effects ambiguous in some domains.

Perhaps the best cases for AI agents reducing economic rents are in markets where firms exploit behavioral biases and bounded rationality. For example, many consumers choose suboptimal contracts given their expected usage patterns, leading to rents for firms (e.g., see Grubb (2009) on phone contracts). By making more rational decisions on behalf of consumers, agents would render such rent-extraction strategies less viable, pushing markets closer to the competitive ideal.

Agents also dramatically lower search costs. Search-theoretic models have long emphasized how cognitive and time-based costs prevent consumers from gathering and processing information necessary for optimal purchasing decisions. By rapidly collecting, analyzing, and comparing product data across markets, agents can mitigate these frictions and enhance allocative efficiency.

Beyond immediate efficiency gains, agents also affect longer-term market dynamics. As they consistently direct demand toward higher-quality or better-value offerings, producers would receive clearer market signals about consumer preferences. This feedback mechanism would incentivize firms to invest in producing what consumers want rather than what they can be persuaded to buy through marketing or by exploiting cognitive limitations. It would also incentivize firms to allow for more customization in their products, given that the costs of customization and discovery will drop. The resulting shift in production patterns would compound initial allocative gains, creating a virtuous cycle of better-targeted supply meeting more accurately expressed demand.

AI agents will also affect the prevalence and dynamics of bargaining as their opportunity cost is not human time. Classic bargaining models assume impatience and a preference to conclude negotiations quickly because time and attention are scarce resources for humans. However, for the AI agent the binding constraint is compute rather than time. As a result, even if the principal prefers

timely resolution, agents can initiate negotiations earlier, hold many concurrent negotiations, and continue engaging in these for a longer period of time (e.g., begin negotiating summer-2027 rentals in January 2026).

Beyond bargaining dynamics, AI agents will also alter the prices of goods and services, as agents' ability to elicit and reveal preferences enables new mechanisms.⁷ When consumers willingly and accurately reveal their preferences to trusted agents, firms gain access to richer, more detailed preference information. This enhanced informational access enables personalized pricing strategies that can improve price discovery and market efficiency, lowering the deadweight loss previously caused by information asymmetry. However, efficiency improvements do not determine the distribution of rents. While personalized pricing may benefit consumers by closely matching products to individual preferences, it may also lead to inequitable outcomes if firms exploit consumer-specific price elasticities. At the same time, consumers' agents may strategically withhold certain information to secure more favorable prices, adding additional complexity to the market.

It remains unclear how efficiency improvements will affect the distribution of rents and price dispersion. Many economists initially predicted that the Internet would eliminate price dispersion in markets, reasoning that decreased search costs would create conditions similar to Bertrand competition. However, this prediction largely failed to materialize, perhaps because human search costs suffer from Baumol's cost disease—as productivity increases, so does the opportunity cost of time. A sizable theoretical and empirical literature in economics has investigated persistent price dispersion in the digital era. When products are differentiated, for example, lower search costs can paradoxically lead to higher prices and greater dispersion (Ellison and Ellison 2018; Kaye 2024). AI agents may exacerbate this even further: by better identifying and matching consumer preferences, they can increase effective differentiation and sustain higher dispersion. Furthermore, to counteract the superior capabilities of agents, firms may adopt increasingly sophisticated price obfuscation tactics (Ellison and Ellison 2009).

Another reason why widespread adoption of agents does not guarantee efficient market operation is due to the presence of externalities across agents. For example, reducing the cost of applying to a job with a customized cover letter by deploying AI may flood employers with applications,

⁷AI agents may also act as pricing agents producing similar complications as in Calvano et al. (2020) and Brown and MacKay (2023).

imposing higher screening costs and making it harder to pick the right candidate due to congestion (Wiles and Horton 2025; Kessler 2025). Such market failures highlight the critical role that market design will play in ensuring positive-sum outcomes as AI agents become ubiquitous in economic transactions.

6 Market Design for AI Agents

As AI agents proliferate in economic transactions, markets themselves will need fundamental re-configuration to take advantage of their capabilities and accommodate the unique challenges they offer. This transformation extends beyond policy adjustments to encompass technical infrastructure, identity verification systems, and entirely new market mechanisms that use agents’ superior computational abilities. In this section, we discuss aspects of market design, as well as their application to several important industries.

6.1 Identity

A fundamental challenge of implementing new market designs emerges from the reality that most Internet activity will eventually originate from AI agents rather than humans. These agents will be able to effectively mimic humans in many situations. For example, advertisements are priced based on impressions and clicks, but agents will also be able to “see” and click on them. Similarly, social media companies and their users would like to clearly verify that content is produced by a specific human rather than a bot. Sybil attacks—where single entities create multiple fake identities to gain disproportionate network control—become more difficult to detect in agent-mediated environments. This necessitates new approaches to digital accountability, potentially including cryptographic proofs, digital IDs, or decentralized verification platforms to ensure system safety and reliability.

Of particular importance is identity verification, both for humans and for AI systems that are acting on behalf of specific humans. Two broad classes of solutions emerge: walled-garden approaches where platform gatekeepers exclude potentially malicious actors, and open systems with robust verification mechanisms. In the walled-garden approach, platforms require a log-in prior to interacting with content. This approach is imperfect. A person can launch an agent after

logging-in, creating spam and malicious content. If the platform bans this person, this same person can create additional accounts.

Another approach is a proof-of-personhood system, which tries to create a network of unique humans. For example, the World Foundation has created a biometric technology that uses the human iris to uniquely identify humans and to give them access to an app that can be used to prove this uniqueness (World Foundation 2023). Proof of humanhood solutions can prevent sybil attacks and can preserve privacy, but require a larger restructuring of existing systems and mass adoption to be truly transformative.

Identity systems will also need to be combined with verified credentials and reputation mechanisms. For example, consider an advertiser wishing to target humans with particular demographics. Some of these demographics are verifiable using a digital government identity, and this identity can be used to reveal demographics to a particular advertiser (European Commission 2025). Reputation mechanisms can also be designed to be attached to particular identities in a similar manner. Combined, identity, credentials, and reputation can enable more sophisticated market interactions, including targeted deals and negotiations over small purchases that would previously have been economically infeasible.

6.2 Changes in Existing Platform Designs

Platforms of all types will need to change in response to the growth of AI agents. We speculate about the types of changes that might occur through three examples.

First, consider the fact that agents may consume content such as social media posts and search results in e-commerce. AI will serve as a filter between humans and the broader digital environment. Much of today’s online content combines user-desired material with content designed primarily to capture attention or generate revenue, from native advertising to engagement-optimized recommendation algorithms. AI agents could review incoming information streams and selectively transmit content based on user preferences and actual utility, substantially reducing exposure to irrelevant material and ultimately steering digital platforms toward prioritizing user preferences over attention capture. This filtering function poses challenges to existing digital business models that rely on bundling content with advertising or engagement-driven recommendations.

Firms often create choice architectures that lead users to act in ways beneficial to the firm.

Rational AI agents browsing on behalf of users are unlikely to be influenced by these nudges. Similarly, advertising is often characterized by puffery, and selective information disclosure. AI agents may be much less influenced by these types of ads than humans. As agents proliferate, platforms will need to adjust. Platforms will try to design information architectures that specifically influence agents, something that is already evident with the nascent field of optimization of content for Large Language Model (LLM) consumption. Over time, we may see a shift to alternative platform monetization models, such as subscription services, which have less dependence on advertising.

Next, consider AI agents as content creators. Agents will be able to create and like social media posts, create artistic works, send direct messages, offer to buy items, and apply to jobs. Since the costs of these actions are low, it may be individually rational for humans to ask their agents to do these actions en masse (Goldberg and Lam 2025). Yet part of the value of these actions is that they serve as signals. Merely the fact that a piece of content was produced by a human versus an AI may be important for signaling and consumption utility (Longoni et al. 2022; Rae 2024). In response to content proliferation, platforms will be forced to adapt rules restoring costly signaling and credible verification that content was created by a human. For example, posters on social media may be required to pay per post, so that the platform does not get flooded with low-quality or undifferentiated content. In e-commerce, fraud is already a first-order concern due to fake listings and credit card chargebacks. We expect platforms to add monetary costs or other limits on these activities. We have already examined how digital identity might function in a world populated by AI agents. Platforms must now determine which forms of identity, if any, to require, and how to incorporate identity-related information into their operations.

Lastly, infrastructure will also likely change as a result of increased agent usage. For example, currently, users do not pay websites per HTTP request, both because the marginal cost of an individual request is low and because of technical challenges. However, as agentic browsing becomes ubiquitous, overall traffic volumes are likely to surge, leaving website owners to bear these costs. In response, Cloudflare, a content delivery network provider that helps domains manage traffic, recently introduced “pay-per-crawl,” a feature that allows website owners to charge agents for crawling their sites (Allen and Newton 2025). This new capability for websites reflects a Pigouvian logic: market participants should pay for the externalities they impose on others. The new traffic and cost profile may require entirely new, agent-first surfaces—authenticated, rate-limited APIs

with machine-readable pricing and consent signals—rather than human-oriented pages. We expect new conventions to emerge, akin to the robots.txt standard, that define how agents interact with platforms and what activities are permitted or prohibited. Over time, technical evolution may also sharpen the distinction between agent-oriented and person-oriented interfaces, with websites offering parallel access: streamlined, data-rich endpoints optimized for agents alongside traditional interfaces for humans.

6.3 Enabling Previously Impractical Market Designs

Perhaps the most consequential and transformative change AI agents enable is the practical deployment of theoretically superior market designs that have long remained underutilized. Consider the deferred acceptance algorithm (Gale and Shapley 1962), which guarantees stability and strategy-proofness for proposing parties. Although many labor markets and matching services could benefit from such mechanisms, their use has been limited because they require comprehensive preference rankings—information that is cognitively demanding for individuals to generate and costly to collect at scale. As a result, most matching platforms (e.g., Upwork, Tinder) rely on far simpler recommendation systems that present all participants with the same “best” options, often leaving outcomes unstable or leaving many participants unmatched.

AI agents can significantly alter matching markets. Trained to reason directly over natural language, foundation models already allow a single agent to parse a paragraph describing one’s tastes (Rusak et al. 2025). Soon, they will move beyond just facilitating preference elicitation to discovering them. Like a skilled coach or therapist, an AI agent can help users articulate values by detecting behavioral patterns: a home-buying assistant might notice a buyer consistently favoring houses near parks with large windows—suggesting unspoken preferences for green space and natural light. By surfacing these insights and prompting reflection, agents could help individuals better understand their own priorities. With AI-derived preferences, markets could implement sophisticated matching algorithms requiring preferences over thousands of alternatives. Labor markets could have both job seekers and employers provide natural language descriptions of their preferences, enabling comprehensive rankings that support equilibrium matching algorithms superior to traditional recommendation systems (Rusak et al. 2025). This capability extends beyond labor markets to any matching context where the complexity of optimal allocation exceeds human computational

capacity.

Agents also enable new designs for privacy and strategic transparency. By delegating sensitive questions to AI agents, parties can credibly commit to privacy-preserving interactions. Both sides can precommit their agents to pose all legally permissible questions without penalty, with agents disclosing only relevant responses. These mechanisms may effectively separate sensitive inquiries from signaling concerns, enhancing transparency and efficiency across sensitive contexts. For example, a job seeker may hesitate to directly inquire about maternity leave policies to avoid negative signaling to employers, despite having a legitimate interest in obtaining this information. AI intermediaries can resolve these tensions by posing sensitive queries on behalf of users within precommitted, privacy-preserving protocols. As another example, the classic case of Disney’s anonymous land purchases for Walt Disney World demonstrates how strategic anonymity can prevent price manipulation. However, market designers must carefully balance anonymity benefits against verification needs, considering whether multiple personas linked to single individuals should be permitted (Buterin 2025). Taken together, AI agents enable new market designs that have previously unimaginable benefits.

7 Regulation

There are already broad AI-related policy efforts such as municipal ordinances, federal initiatives, and technical safeguards. As agents become increasingly integrated into decision-making processes, regulatory frameworks will need to evolve to address challenges around market power, autonomy, liability, privacy, and intellectual property rights. While we briefly examine each of these areas, the discussion is illustrative rather than exhaustive.

1. Market Power: A central concern in AI regulation is the concentration of market power among a few large firms that control the compute, data, and distribution needed to build frontier models, creating high barriers to entry. For consumers, this can mean agents that are less independent or customizable and biased toward platform interests. Reduced competition also limits diversity, constrains interoperability, and weakens users’ ability to carry agents across contexts—resulting in fewer choices, weaker alignment, and less control. Regulation must therefore protect

consumer autonomy by ensuring access to essential infrastructure, mandating interoperability, and preventing exclusionary practices, while guarding against regulatory capture that entrenches incumbents and leaves consumers with diminished functionality and freedom.

Current developments illustrate the stakes. As of September 2025, antitrust scrutiny of large AI firms has intensified, including investigations into Google’s market practices (U.S. Department of Justice 2025). These cases highlight the delicate balance regulators face: while underregulation risks leaving unchecked monopolistic behavior, overregulation could inadvertently stifle technological innovation and slow beneficial applications of AI.

2. Autonomy and Liability: With AI agents acting on behalf of humans, the question of who bears responsibility when things go wrong becomes unavoidable. Whether accountability should rest with the human who delegates the final decision to an AI system or the human who acts on another party’s AI-generated output (i.e., under negligence liability), or with the firms that develop or deploy such systems regardless of fault (i.e., strict liability), remains contested. This choice raises fundamental concerns about how to allocate the costs of harm in a way that protects users without choking off innovation.

The European Union’s (EU) recent product liability directive marks the most advanced regulatory attempt to grapple with these challenges to date, explicitly extending liability to digital goods, software, and ongoing updates (European Union 2024). By recognizing that the traditional boundary between a “finished product” and its subsequent behavior no longer holds in the age of AI—as these systems learn, adapt, and sometimes act in ways that surprise even their creators—the directive exemplifies how adjustments to existing tort frameworks, rather than creation of entirely new ones, may become the central lever shaping the trajectory of AI development and adoption.

3. Security and Privacy: Even with clearer ex-post liability measures, AI systems remain vulnerable to adversarial manipulation. Jailbreaking attacks, for example, allow users to circumvent safety measures and access system features in unintended ways, representing a persistent challenge for current chatbot technology (Shen et al. 2024). The vulnerabilities become particularly concerning in high-stakes applications like hiring, where malicious actors might exploit AI interviewers through carefully crafted prompts to gain unfair advantages.

These failures raise privacy concerns: once guardrails are bypassed, agents can pull data across contexts, retain it beyond its intended use, or infer it from seemingly benign traces. Training and adaptation on sensitive user data, often without explicit consent, heightens these risks and raises questions about data ownership and control, especially in employment, housing, and other high-stakes domains. Alternatively, overly restrictive privacy regulation may prevent agents from using private data for the purposes of better alignment.

Another privacy risk stems from inadvertent training on private or sensitive data. For example, a model trained based on Bob’s social media posts may memorize information about Bob. When Amy writes a prompt relating to Bob, information about Bob may surface in unpredictable and potentially damaging ways. In addition to the leakage of private information memorized in training, another risk may simply be sophisticated inference from usage traces like late-night browsing or vitamin purchases. This information can be repurposed in other contexts (e.g., lower salary offers based on perceived pregnancy status or higher insurance premiums based on inferred health risks). Due to these concerns, we expect that existing data regulations (California Consumer Privacy Act 2018; General Data Protection Regulation 2018) will need to be adapted in response to generative AI and automated decision-making.

4. Data Rights and Platform Access: Generative AI models, and thus agents, are trained on vast repositories of creative content, much of it scraped without authorization. Once such data has been incorporated, “unlearning” is technically difficult, raising questions about compensation for artists, writers, and other creators. But beyond training data, a parallel issue arises in deployment: how external agents access and use live platform data.

Platforms are already taking actions to block agents from accessing their content via litigation (The New York Times Co. v. Microsoft Corp. 2023), robots.txt rules (Smith 2025), and tighter terms of service. Yet more capable agents can emulate human browsing or operate on the principal’s device to conduct certain actions. This creates a distinctive tension: from the platform’s perspective, third-party agents extract value by crawling data and intermediating transactions without bearing the costs of producing, curating or securing that data. From the agent provider’s perspective, platforms’ restrictions can look like attempts to monopolize access and limit competition.

The central concern is thus not just whether external agents should have access, but on what

terms—in particular, whether platforms will be required to license or otherwise compensate for data usage when agents intermediate user activity. Without clear mechanisms for compensation, platforms may under-invest in data quality, while overly restrictive access could stifle competition and consumer choice. Resolving these concerns requires additional litigation and regulation.

8 Conclusion

The capacity of AI agents to dramatically reduce transaction costs as automated intermediaries could unlock new forms of market participation, enable previously infeasible mechanisms, and push allocative efficiency closer to competitive ideals. Yet the same forces that make agents attractive—their tireless persistence, computational superiority, and negligible marginal costs—also threaten to overwhelm existing market structures. The ultimate impact will depend critically on collective choices adopted regarding agent design, market structures, and regulatory frameworks.

While the capabilities of AI agents are new, their transformative impact is ultimately governed by the fundamental economic forces such as supply and demand. In the end, these forces will decide whether agents become a foundation for welfare-enhancing prosperity or collapse into merely rent-redistributive outcomes. Economists have the tools to play a key role in this transformation, but whether they choose to do so remains to be seen.

Acknowledgments

We thank David Rothschild and Anton Korinek for their detailed comments and feedback.

References

- Allen, Will, and Simon Newton. 2025. “Introducing Pay Per Crawl: Enabling Content Owners to Charge AI Crawlers for Access.” *Cloudflare Blog*. Accessed September 29, 2025.
- Allouah, Amine, Omar Besbes, Josué D. Figueroa, Yash Kanoria, and Akshit Kumar. 2025. “What Is Your AI Agent Buying? Evaluation, Implications and Emerging Questions for Agentic E-Commerce.” *arXiv preprint arXiv:2508.02630*.
- Bostrom, Nick. 2014. *Superintelligence: Paths, Dangers, Strategies*. Oxford: Oxford University Press.

- Brown, Zach Y., and Alexander MacKay. 2023. “Competition in Pricing Algorithms.” *American Economic Journal: Microeconomics* 15 (2): 109–56.
- Buterin, Vitalik. 2025. “Does digital ID have risks even if it’s ZK-wrapped?” June 28. Accessed September 15, 2025.
- California Consumer Privacy Act of 2018. Cal. Civ. Code § 1798.100 et seq., as amended. “California Consumer Privacy Act (CCPA).” California Attorney General. Updated March 13, 2024. Accessed September 15, 2025.
- Calvano, Emilio, Giacomo Calzolari, Vincenzo Denicolo, and Sergio Pastorello. 2020. “Artificial Intelligence, Algorithmic Pricing, and Collusion.” *American Economic Review* 110 (10): 3267–97.
- Christian, Brian. 2020. *The Alignment Problem: Machine Learning and Human Values*. New York: W. W. Norton & Company.
- Citron, Dave. 2024. “Try Deep Research and Our New Experimental Model in Gemini, Your AI Assistant.” *Google Blog*, December 11. <https://blog.google/products/gemini/google-gemini-deep-research/>.
- Coase, Ronald H. 1937. “The Nature of the Firm.” *Economica* 4: 386–405.
- Dammu, Preetam, Prabhu Srikar, Omar Alonso, and Barbara Poblete. 2025. “A Shopping Agent for Addressing Subjective Product Needs.” In *Proceedings of the Eighteenth ACM International Conference on Web Search and Data Mining (WSDM ’25)*, 1032–35.
- Ellison, Glenn, and Sara Fisher Ellison. 2009. “Search, Obfuscation, and Price Elasticities on the Internet.” *Econometrica* 77 (2): 427–52.
- Ellison, Glenn, and Sara Fisher Ellison. 2018. “Match Quality, Search, and the Internet Market for Used Books.” NBER Working Paper w24197. Cambridge, MA: National Bureau of Economic Research.
- European Commission. 2025. *European Digital Identity*. Accessed September 26, 2025. https://commission.europa.eu/strategy-and-policy/priorities-2019-2024/europe-fit-digital-age/european-digital-identity_en.
- European Union. 2024. Directive (EU) 2024/2853 of the European Parliament and of the Council of 23 October 2024 on liability for defective products and repealing Council Directive 85/374/EEC. OJ L 2024/2853, 18 November 2024. Accessed September 29, 2025. <http://data.europa.eu/eli/dir/2024/2853/oj>.
- Fradkin, Andrey. 2025. “Demand for LLMs: Descriptive Evidence on Substitution, Market Expansion, and Multihoming.” *arXiv preprint arXiv:2504.15440*.
- Gale, David, and Lloyd S. Shapley. 1962. “College Admissions and the Stability of Marriage.” *The American Mathematical Monthly* 69 (1): 9–15.
- General Data Protection Regulation (GDPR). 2016. Eur. Parl. & Council Regulation 2016/679, OJ L 119, 4 May 2016; corrected in OJ L 127, 23 May 2018.
- Goldberg, Samuel, and H. Tai Lam. 2025. “Generative AI in Equilibrium: Evidence from a Creative Goods Marketplace.” SSRN Working Paper 5152649.

- Grubb, Michael D. 2009. “Selling to Overconfident Consumers.” *American Economic Review* 99 (5): 1770–807.
- Hadfield, Gillian K., and Andrew Koh. 2025. “An Economy of AI Agents.” *NBER The Economics of Transformative AI*. https://conference.nber.org/conf_papers/f227497.pdf.
- Jabarian, Brian, and Luca Henkel. 2025. “Voice AI in Firms: A Natural Field Experiment on Automated Job Interviews.” SSRN Working Paper 5395709.
- Kalai, Adam Tauman, Ofir Nachum, Santosh S. Vempala, and Edwin Zhang. 2025. “Why Language Models Hallucinate.” *arXiv preprint arXiv:2509.04664*.
- Kaye, Aaron. 2024. “The Personalization Paradox: Welfare Effects of Personalized Recommendations in Two-Sided Digital Markets.” Working Paper.
- Kessler, Sarah. 2025. “Employers Are Buried in A.I.-Generated Résumés.” *New York Times*, June 21. Accessed September 15, 2025.
- Liang, Annie. 2025. “Artificial Intelligence Clones.” *arXiv preprint arXiv:2501.16996*.
- Longoni, Chiara, Andrey Fradkin, Luca Cian, and Gordon Pennycook. 2022. “News from Generative Artificial Intelligence Is Believed Less.” In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*, 97–106.
- The New York Times Co. v. Microsoft Corp. 2023. No. 1:23-cv-11195 (S.D.N.Y., December 27). Complaint. Accessed September 9, 2025. https://nytco-assets.nytimes.com/2023/12/NYT_Complaint_Dec2023.pdf.
- Rae, Irene. 2024. “The Effects of Perceived AI Use on Content Perceptions.” In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*.
- Rothschild, David M., Markus Mobius, Jake M. Hofman, Eleanor W. Dillon, Daniel G. Goldstein, Nicole Immorlica, Sonia Jaffe, Brendan Lucier, Aleksandrs Slivkins, and Matthew Vogel. 2025. “The Agentic Economy.” *arXiv preprint arXiv:2505.15799*.
- Rusak, Gili, Benjamin S. Manning, and John J. Horton. 2025. “AI Agents Can Enable Superior Market Designs.” Working Paper.
- Shen, Xinyue, Zeyuan Chen, Michael Backes, Yun Shen, and Yang Zhang. 2024. “Do Anything Now: Characterizing and Evaluating In-the-Wild Jailbreak Prompts on Large Language Models.” In *Proceedings of the 2024 ACM SIGSAC Conference on Computer and Communications Security*.
- Smith, Allison. 2025. “Shopify Quietly Sets Boundaries for AI Agents on Merchant Sites.” *Modern Retail*, July 14, 2025. Accessed September 30, 2025. <https://www.modernretail.co/technology/shopify-has-quietly-set-boundaries-for-buy-for-me-ai-bots-on-merchant-sites/>.
- U.S. Department of Justice, Office of Public Affairs. 2025. “Department of Justice Prevails in Landmark Antitrust Case Against Google.” April 17. Accessed September 29, 2025. <https://www.justice.gov/opa/pr/departments-justice-prevails-landmark-antitrust-case-against-google>.
- Wiles, Emma, Zanele Munyikwa, and John J. Horton. 2025. “Algorithmic Writing Assistance on Jobseekers’ Resumes Increases Hires.” *Management Science*.

- Wiles, Emma, and John J. Horton. 2025. “Generative AI and labor market matching efficiency.” SSRN Working Paper 5187344.
- World Foundation. 2023. “Iris Recognition Inference System.” World.org, June 5. <https://world.org/blog/engineering/iris-recognition-inference-system>.
- Zeff, Maxwell. 2025. “Google Rolls Out Project Mariner, Its Web-Browsing AI Agent.” *TechCrunch*, May 20. <https://techcrunch.com/2025/05/20/google-rolls-out-project-mariner-its-web-browsing-ai-agent/>.
- Zhu, Shenzhe, Jiao Sun, Yi Nian, Tobin South, Alex Pentland, and Jiaxin Pei. 2025. “The Automated but Risky Game: Modeling Agent-to-Agent Negotiations and Transactions in Consumer Markets.” *arXiv preprint arXiv:2506.00073*.