

NBER WORKING PAPER SERIES

WHO WANTS TO BREAK UP BIG FIRMS? HARM, FAIRNESS, AND THE DEMAND FOR
ANTITRUST

Ricardo Perez-Truglia
Jeffrey Yusof

Working Paper 35503
<http://www.nber.org/papers/w35503>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
July 2026, Revised August 2026

This draft: August 7, 2026. We thank Alexia Witthaus Vine for excellent research assistance. We are thankful for comments by colleagues. We gratefully acknowledge funding from UCLA's Anderson Behavioral Lab. The experiment was pre-registered at the AEA RCT Registry (AEARCTR-0018249) and was reviewed and approved in advance by the North General Institutional Review Board at UCLA (IRB-26-0751). After the study is accepted for publication, we will share all code and data through a public repository. The views expressed herein are those of the authors and do not necessarily reflect the views of the National Bureau of Economic Research.

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 They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2026 by Ricardo Perez-Truglia and Jeffrey Yusof. 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.

Who Wants to Break Up Big Firms? Harm, Fairness, and the Demand for Antitrust
Ricardo Perez-Truglia and Jeffrey Yusof
NBER Working Paper No. 35503
July 2026, Revised August 2026
JEL No. C90, D83, K21, L40

ABSTRACT

The rise of superstar firms has made dominant companies central to modern economic life, and antitrust enforcement is one of the main policy tools for regulating their market power. Public opinion can shape the political and regulatory environment in which antitrust enforcement takes place, yet there is little direct evidence on what drives these preferences. We conduct a pre-registered information-provision experiment with 4,000 American households. Respondents were told about one of five real antitrust cases and randomly assigned to information treatments designed to study four potential drivers of support for antitrust enforcement: perceived market share, perceived consumer harm, perceived unfair competition, and perceived negative image. All four treatments moved the beliefs they were designed to affect, but their effects on demand for antitrust differed sharply. Information about consumer harm had the most systematic effects: it increased plaintiff support and support for break-up and conduct remedies, with effects remaining visible one month later, and also spilled over to broader support for antitrust policies. By contrast, and contrary to expert forecasts, information about market share had no meaningful effect on demand for antitrust enforcement. The findings suggest that the public thinks like economists in one key respect: they do not care about market share per se, but respond instead to consumer harm. One factor outside the core economic framework, perceived unfair competition, also matters, though its effects are more limited in scope. We discuss implications for policymakers and regulators.

Ricardo Perez-Truglia
University of California, Los Angeles
and NBER
ricardotruglia@gmail.com

Jeffrey Yusof
University of Stuttgart
jeffrey.yusof@ivr.uni-stuttgart.de

An online appendix is available at <http://www.nber.org/data-appendix/w35503>
A randomized controlled trials registry entry is available at
<https://www.socialscienceregistry.org/trials/18249>

1 Introduction

Superstar firms have become a defining feature of the U.S. economy as scalable markets, winner-take-all dynamics, and network externalities increasingly concentrate economic activity among a small number of dominant companies. For example, [Autor et al. \(2020\)](#) document rising sales concentration from 1982 to 2012 across all major sectors. This concentration of power is salient in capital markets too: according to Dow Jones Market Data, the ten largest companies in the S&P 500 represented 43.2% of the index’s total market value in June 2026 ([Wall Street Journal, 2026](#)). Their significance extends beyond product and capital markets, as these firms and their leaders can also shape the political environment, among other things, through lobbying and campaign spending ([Perez-Truglia and Yusof, 2026](#)).¹

Dominant firms can generate large benefits for consumers through scale, convenience, and innovation. At the same time, their market dominance creates the risk of classic antitrust harms: higher prices, lower quality, reduced innovation, exclusion of rivals, and fewer choices for consumers ([Shapiro, 2019](#)). As a result, antitrust enforcement has re-emerged as a central policy tool for regulating dominant firms and preserving competition ([New York Times, 2024](#)). The remedies at stake can be substantial, ranging from restrictions on specific business practices to divestitures or breakups, and shifts in antitrust enforcement have the potential for large effects on firm values and market outcomes ([Baker et al., 2023](#)).

Although antitrust enforcement is often described as a technical legal and economic domain, public opinion still matters. In some cases, it can matter directly: private damages actions may be decided by lay juries rather than only by judges or regulators.² More importantly, public opinion can shape antitrust enforcement indirectly by affecting the political incentives of elected officials who write antitrust laws, fund enforcement agencies, and oversee the broader direction of competition policy. Presidents also appoint the leaders of the Federal Trade Commission and the Department of Justice, including officials who may take a more or less aggressive approach to antitrust enforcement ([Washington Post, 2024b](#)). Because these elected officials have incentives to respond to their constituents, whether they pursue more aggressive antitrust policies may depend at least in part on what the public demands, a link emphasized by historical work on the political economy of U.S. antitrust enforcement ([Lancieri et al., 2023](#)).

¹For example, large tech companies have spent substantial amounts on federal lobbying ([Axios, 2026](#); [Fortune, 2026](#)), and billionaires with major stakes in these companies have become increasingly influential through their campaign contributions and ownership of media outlets ([Washington Post, 2024a](#)).

²The right to a jury trial applies to treble-damages suits under the antitrust laws; see [U.S. Supreme Court \(1959\)](#). A prominent recent example is *Epic Games v. Google*, in which a federal jury found that Google violated federal and California antitrust laws in connection with the Google Play Store; the Ninth Circuit later affirmed the jury verdict and injunction ([U.S. Court of Appeals for the Ninth Circuit, 2025](#)).

Despite the importance of public opinion for antitrust enforcement, there is limited direct evidence on what drives public demand for it. We fill this gap with a pre-registered survey experiment studying Americans’ preferences over antitrust enforcement and the mechanisms underlying those preferences.

Our design is organized around four beliefs that we hypothesize may shape demand for antitrust enforcement. The standard economic framework evaluates antitrust enforcement through its effects on welfare (Tirole, 1988; Shapiro, 2019). In this view, the central question is not whether a firm has a large market share or faces few competitors, but whether market power harms consumers, for example, through higher prices. This motivates the first belief, *perceived consumer harm*. The other three beliefs capture considerations that may matter to the public even if they fall outside the standard welfare-based framework. Respondents may react to *perceived market share* itself. A large market share is not necessarily harmful, since it may reflect a product or service that consumers genuinely prefer, greater convenience, or other non-exclusionary advantages (Demsetz, 1973), but the public may fail to reason like economists. The public may also care about *perceived unfair competition*. Even for the same market outcome, respondents may react differently depending on whether they see it as the result of fair competition or of conduct they perceive as unfair, such as exclusivity agreements that make it harder for rivals to compete. Finally, broader moral judgments about the firm may affect the public’s demand for enforcement even when the information is not directly related to competition, which we capture with the company’s *perceived negative image*. Because these beliefs are conceptually distinct but empirically interdependent, we require a research design that can disentangle the causal effect of each belief even when multiple beliefs may move together.

Subjects are randomly assigned to one of five real antitrust cases: Google vs. the Department of Justice, Meta vs. the Federal Trade Commission, Live Nation vs. the Department of Justice, Apple vs. Epic Games, or a class action lawsuit against Luxottica.³ We focus on real cases because they make the information more concrete and credible, and because they allow us to study demand for antitrust enforcement in actual policy disputes rather than in abstract scenarios. After introducing respondents to their assigned case, we elicit beliefs about the four dimensions discussed above: perceived market share, perceived consumer harm, perceived unfair competition, and perceived negative image.

Respondents are then randomly assigned to one of five information conditions. The control group receives no additional information. The four treatment groups receive fact-based information about one of the four belief dimensions: market share, consumer harm,

³Essilor and Luxottica combined in 2018 to create EssilorLuxottica. For brevity, in the survey materials and in the paper we refer to the company simply as Luxottica.

unfair competition, or negative image. The treatments follow the same conceptual structure across cases, but the specific content is tailored to the assigned case. They provide short arguments and facts supported by documentation such as government reports, academic research, expert surveys, newspaper articles, and excerpts from television news coverage. For the sake of brevity, we use the Google case as the running example in the rest of the introduction, but the same design was implemented for all five cases. In the Google case, for example, the treatments provided information about Google’s share of online search advertising, how advertising costs may be passed through to consumers, its payments to be the default search engine on browsers and devices, or a controversy involving the company that was unrelated to antitrust.

After the information-provision stage, we re-elicited the same four beliefs, allowing us to measure how each treatment changes respondents’ beliefs. We then measure the main outcomes of interest: support for the plaintiff in the assigned case, support for a structural remedy, and support for a conduct remedy. In the Google case, these outcomes are whether respondents side with the Department of Justice, whether they support breaking up Google, and whether they support restricting Google’s default-search agreements. To measure whether the treatment effects extend beyond the assigned case to views about other dominant firms, we ask analogous belief and outcome questions about dominant firms more broadly. Some of the policy-support measures are behavioral, such as signing a petition, donating to an antitrust NGO, and sending an email to a representative. To measure the persistence of the effects, one month after the baseline survey, we invited respondents to complete a follow-up survey that re-elicited belief and outcome measures.

We conducted the baseline survey in May 2026 with 4,016 American respondents recruited through Prolific, 71.3% of whom also completed the follow-up survey. We first show that the information treatments strongly affected beliefs. The market-share treatment increased perceived market share by 6.8 percentage points (pp), equivalent to 0.39 standard deviations (SD). The consumer-harm treatment increased perceived consumer harm by 0.48 SD, while the unfair-competition and negative-image treatments increased their targeted beliefs by 0.53 SD and 0.42 SD, respectively. These belief effects were not fully compartmentalized. Market-share information primarily shifted perceived market share, with limited spillovers to other beliefs. The consumer-harm, unfair-competition, and negative-image treatments had the strongest effects on their targeted beliefs but also generated more substantial spillovers to other beliefs.

Next, we discuss the average effects of each treatment. We organize the treatment-effect analysis around three dimensions: intensity, breadth, and durability. Intensity asks whether information moves respondents beyond plaintiff support toward support for concrete

remedies, such as breaking up the company; breadth asks whether information about one antitrust case changes views about dominant firms and antitrust policy more generally; and durability asks whether these effects remain visible one month later.

While the market-share treatment has a substantial effect on perceived market share, it lacks intensity—it does not meaningfully increase demand for enforcement or remedies. The consumer-harm treatment has the strongest overall profile: it increases plaintiff support by 0.45 SD and support for remedies by 0.28 SD, generates the clearest spillovers to broader antitrust beliefs and policy support, and has the strongest durability in the follow-up survey. The unfair-competition information increases plaintiff support by 0.48 SD and remedy support by 0.33 SD, but its effects are less broad and less durable. Negative-image information has narrower effects, increasing plaintiff support at baseline but generating little breadth or durability.

The average treatment effects are informative, but they do not by themselves reveal which beliefs causally drive demand for antitrust enforcement. Because each treatment can move multiple beliefs at once, the effect of any given treatment may reflect a combination of changes across the four beliefs. We therefore estimate a 2SLS model that can disentangle the causal effect of each belief. The results point to a clear hierarchy of mechanisms. Perceived consumer harm is the broadest driver of demand for antitrust enforcement. Perceived unfair competition also matters, especially for plaintiff support and remedies in the assigned case. Perceived negative image plays a more limited role, mainly shifting respondents toward the plaintiff side without generating comparable support for more consequential remedies. Perceived market share has no meaningful independent effect.

These findings were not obvious *ex ante*, as shown by the results from an expert forecast survey. To benchmark our findings, we invited a panel of experts on the topic, described the experimental design and information treatments, and asked them to predict the treatment effects. Several key results run counter to the expert forecasts. While experts expected all treatments to increase demand for antitrust, we find that some treatments had positive effects while others did not. Moreover, when asked about the causal mechanisms, experts predicted that perceived market share would be the most important channel, whereas the data show that this belief has a clear null effect.

Taken together, the results suggest that public demand for antitrust aligns closely with the standard economic framework. The alignment comes from the central role of consumer harm: respondents do not demand antitrust action simply because a firm is large, but because they believe its market power harms consumers. Perceived unfairness, a factor outside the standard economic framework, also matters, though with less breadth.

These findings have implications for how policymakers and regulators communicate about

antitrust enforcement. Public discussions often emphasize the size and market shares of dominant firms, but our results suggest that this information is unlikely to increase support on its own. Messages focused on prices, quality, choice, and other consumer-welfare consequences may be more effective. Fairness-based arguments about exclusionary or unfair conduct can also increase support for case-specific remedies, but appear less likely to generate broader support for antitrust policy.

This paper relates to and contributes to two strands of literature. First, it contributes to the literature on attitudes toward inequality and preferences for redistribution. A large literature studies how the public responds to the concentration of income and wealth among individuals, including preferences for taxing billionaires, business elites, and the top 1% (Di Tella et al., 2021; Hope et al., 2023; Perez-Truglia and Yusof, 2026). We contribute to this literature by shifting attention from the concentration of resources among individuals to concentration among firms. This distinction is important because the concentration of economic resources is not limited to individuals: firms also differ dramatically in size and market power, yet we know much less about how the public responds to those disparities.

Second, this study relates to work on market power and antitrust. A large literature in economics and law studies how market power affects prices, quality, innovation, entry, and welfare, as well as how antitrust law should define markets and evaluate competitive harm (Tirole, 1988; Shapiro, 2019). Despite its importance, no economics study directly measures what drives demand for antitrust enforcement.⁴ We contribute to this literature by filling that gap. We show that Americans' preferences align with the economic framework underlying this literature: they do not care about *market share* per se, but care about *market power* instead.

The rest of the paper proceeds as follows. Section 2 describes the research design and implementation of the experiment. Section 3 presents descriptive evidence on baseline demand for antitrust enforcement. Section 4 reports the average treatment effects on beliefs and outcomes. Section 5 studies mechanisms using belief updating and a 2SLS framework. Section 6 concludes.

⁴The closest exceptions are from political science: Brutger and Pond (2025a,b) use survey experiments to study how factors such as partisanship, international competition, and nationalism shape support for strengthening antitrust laws.

2 Research Design

2.1 Hypotheses

This subsection presents the paper’s main hypotheses about the beliefs that may shape demand for antitrust enforcement. We discuss the intuition behind each of the four beliefs below. For a more formal analysis, Appendix A provides a stylized model of market power and incumbent advantage that formalizes this discussion.

The **perceived consumer harm** belief captures the standard economic rationale for antitrust enforcement. The standard economic framework evaluates antitrust enforcement through its effects on welfare (Tirole, 1988; Shapiro, 2019). In this view, the number of competitors or the market shares of individual firms are not the ultimate objects of interest. Instead, the central question is whether market power harms consumers through higher prices, lower quality, fewer options, or less innovation. Thus, if respondents think like economists, the key belief that should drive demand for antitrust is whether the company’s market power produces consumer harm, for example through higher prices.

The **perceived market share** belief captures respondents’ views about the company’s market share. Unlike consumer harm, a high market share is not necessarily harmful. A firm may have a large share because it offers a product or service that consumers genuinely prefer, greater convenience, or other non-exclusionary advantages (Demsetz, 1973). According to the economic framework, holding consumer harm constant, market share should therefore be irrelevant for demand for antitrust. However, to the extent that the public lacks economic training or faces cognitive limitations, it may fail to understand that market share is not harmful per se. Evidence from other policy domains, such as taxation, shows that public reasoning can diverge from economists’ reasoning (Sapienza and Zingales, 2013; Stantcheva, 2021); there is also evidence that members of the public prefer certain policies that economists view as inefficient relative to available alternatives (Jiang et al., 2026).

The **perceived unfair competition** belief captures views about the fairness of the sources of market outcomes. One recurrent topic in antitrust cases is that firms may create barriers that make it harder for consumers, business partners, or rivals to move away from the incumbent, including through defaults, exclusive arrangements, or other forms of lock-in (Klemperer, 1987). Such actions could increase market power and consumer harm, but the public may also perceive them as unfair. This perception of unfairness may matter for demand for antitrust. Evidence from survey experiments suggests that fairness norms constrain public acceptance of firms’ profit-seeking behavior (Kahneman et al., 1986). Furthermore, evidence from taxation shows that fairness considerations also shape policy support, such as stronger support for wealth taxes when wealth is perceived as inherited rather than earned (Fisman

et al., 2020; Stantcheva, 2021). Holding market outcomes fixed, such as consumer harm or market share, the public may care about how those outcomes came about. Respondents may demand antitrust if they believe consumer harm arises because consumers do not shop around for the best prices. But if they believe the same consumer harm comes from unfair actions by firms, they may react even more strongly.

The **perceived negative image** belief captures broader moral judgments about the company that may matter for public demand even when they are not part of the standard welfare-based rationale for antitrust enforcement. Intuitively, given a certain level of consumer harm, the public may demand more antitrust if they view the firm as morally objectionable. According to this hypothesis, information related to the morality of the firm may affect demand for antitrust enforcement even if the information is unrelated to competition, for example sexual misconduct by senior executives or gender-based pay discrimination. This hypothesis is inspired by related work showing that demand for taxing the rich can increase when the public learns about their lavish lifestyles (Perez-Truglia and Yusof, 2026) or becomes more distrustful of business elites (Di Tella et al., 2021).

These four beliefs are conceptually distinct but empirically interdependent. A single piece of information can easily shift multiple beliefs at the same time. Consider, for example, information that a firm used exclusivity agreements that made it harder for rivals to compete. This information may increase perceived unfairness, but it can also affect the other beliefs. Because the same conduct can protect the firm’s market position and soften competitive pressure, it may increase the perceived market share and the perceived consumer harm. To the extent that the conduct is seen as revealing something broader about the company, this information may also harm its overall image. This interdependence motivates a research design that measures each belief before and after the information treatments, and an empirical framework that can help disentangle the role of each belief even when the beliefs move together.

2.2 Overview of the Experimental Design

The baseline survey instrument is attached as Appendix F and described in detail below.⁵ Panel A of Figure 1 summarizes the structure of the baseline survey. Respondents were randomly assigned to one of five antitrust cases. After a brief introduction to the assigned case, we elicited four beliefs about that case. Respondents were then randomly assigned to the control group or to one of four information treatments. After the information treatment, we re-elicited the four beliefs, measured the main outcomes, and elicited analogous beliefs

⁵We corrected a handful of minor typographical errors in the survey instruments.

and outcomes about dominant firms more broadly.

2.3 Case Assignment

At the start of the survey, each respondent was randomly assigned with equal probability to one of five antitrust cases.⁶ Each case involved a company alleged to hold a dominant position in a relevant market. Respondents were shown a small number of screens introducing the assigned case. We also included questions assessing respondents' familiarity with the assigned company and antitrust case.⁷ We briefly summarize the five cases below:

- *Google vs. Department of Justice.* The U.S. Department of Justice, joined by state attorneys general, filed an antitrust lawsuit against Google in 2020.⁸ Respondents were told that the case alleged monopolization of internet search and search advertising through exclusionary contracts, especially agreements making Google the default search engine on browsers and devices. They were also told that Google countered that its dominance reflects product quality and consumer choice, and that default-search agreements are standard distribution deals.
- *Live Nation vs. Department of Justice.* The Department of Justice and state attorneys general filed an antitrust lawsuit against Live Nation, the parent company of Ticketmaster, in 2024.⁹ Respondents were told that the case alleged monopoly power across concert promotion, ticketing services, and venue operations, supported by exclusive agreements, coercive tactics, and long-term contracts that limited rival ticketing firms' access to major events. They were also told that Live Nation denied holding a monopoly and argued that ticket prices are largely driven by demand and venue capacity.
- *Apple vs. Epic Games.* Epic Games sued Apple in 2020 over Apple's App Store rules.¹⁰ Respondents were told that Epic alleged that the App Store is an illegal monopoly

⁶The enforcement actor varies across cases: Google, Meta, and Live Nation involve government antitrust enforcement; Apple involves a dispute with Epic Games; and Luxottica involves a class action lawsuit.

⁷Some questions were conditional on earlier responses. For example, respondents who reported being customers or users of the assigned company were asked about their experience with that company.

⁸The DOJ filed the case in October 2020. In August 2024, the U.S. District Court for the District of Columbia concluded that Google had unlawfully maintained monopolies in search and search advertising; in September 2025, the court ordered remedies restricting exclusive distribution contracts and requiring data and syndication access for rivals.

⁹The DOJ filed the lawsuit on May 23, 2024. In March 2026, the DOJ reached a tentative settlement with Live Nation. Some states continued litigating, and in April 2026 a federal jury found that Live Nation and Ticketmaster had operated as an illegal monopoly.

¹⁰Epic filed its lawsuit in August 2020. The district court issued a mixed ruling in September 2021: Apple largely prevailed on Epic's antitrust claims, but the court enjoined Apple's anti-steering rules. In January 2024, the U.S. Supreme Court declined to hear appeals by either side, leaving the lower-court rulings in place.

because Apple controls app installation on iPhones and requires developers to use its in-app payment system, allowing commissions of up to 30 percent. They were also told that Apple countered that the relevant market is broader than the iPhone and that App Store fees reflect security, privacy, and ease-of-use benefits.

- *Class action lawsuit against Luxottica.* The survey described a class action lawsuit against the Luxottica Group, an eyewear company that owns or controls brands such as Ray-Ban and Oakley and retail chains such as LensCrafters, Sunglass Hut, and Pearle Vision.¹¹ Respondents were told that the case alleged anticompetitive practices that inflated prices, restricted consumer choice, and made it difficult for competitors to enter or remain in the eyewear market. They were also told that Luxottica argued that the eyewear industry is highly competitive and that the complaints failed to plausibly allege exclusionary conduct.
- *Meta vs. Federal Trade Commission.* The Federal Trade Commission, joined by state attorneys general, filed an antitrust lawsuit against Meta, which owns Facebook, Instagram, and WhatsApp, in 2020.¹² Respondents were told that the case alleged that Meta maintained monopoly power in social networking by acquiring emerging rivals, especially Instagram and WhatsApp, and by cutting off developers that posed competitive threats. They were also told that Meta rejected the claims, emphasizing that the FTC had previously reviewed the acquisitions and that Instagram and WhatsApp became successful because of Meta's investment.

In selecting the five cases, we sought broad coverage across several dimensions of contemporary antitrust litigation. The cases span different industries, including software, social media, eyewear, and live entertainment; involve companies that vary in public familiarity; and concern different alleged anticompetitive practices, including exclusionary distribution agreements, acquisitions of nascent rivals, app-store restrictions, vertical integration, and long-term contracting practices. They also vary in the nature of the alleged consumer harm, the type of market at issue, and the remedies sought. At the same time, three of the five cases involve technology companies. This emphasis is intentional: it reflects the contemporary antitrust landscape, in which several of the most prominent U.S. antitrust cases have

¹¹The federal class action complaint we identified was filed in July 2023. Related actions were later consolidated as *In re Eyewear Antitrust Litigation*. In September 2025, the U.S. District Court for the Southern District of New York granted defendants' motion to dismiss the amended complaints with leave to replead.

¹²The FTC filed the lawsuit on December 9, 2020. In November 2025, the U.S. District Court for the District of Columbia ruled in favor of Meta; in January 2026, the FTC filed a notice of appeal.

targeted large technology platforms (Wolfe, 2024). The design also allows us to examine heterogeneity by assigned case and assess whether the main patterns are similar across cases.

2.4 Prior Beliefs

Before the information-provision stage, we elicited the four beliefs, which we refer to as the *prior beliefs*.

Table B.1 reports the exact wording and scales for the belief questions, which are summarized below.

- *Perceived market share*: this question provides a market definition and then uses a slider to elicit the respondent’s guess about the share of that market held by the company in question.¹³ For example, respondents assigned to the Google case were asked the share of the U.S. online search advertising market held by Google, and respondents assigned to Live Nation were asked the share of primary ticketing at major U.S. concert venues held by Live Nation.
- *Perceived consumer harm*: this question elicits the respondent’s agreement with the following statement: “[Company]’s market power harms consumers, for example through higher prices or reduced quality.”¹⁴ Subjects were shown a *five-point agreement scale*: strongly disagree, somewhat disagree, neither agree nor disagree, somewhat agree, and strongly agree. In the analysis, we code this question with values 0–4, where higher values correspond to greater perceived consumer harm.
- *Perceived unfair competition*: this question asked: “Do you believe [company]’s market dominance is mainly the result of fair or unfair competition?” We refer to the response options as the *five-point fairness scale*: entirely fair competition, somewhat fair competition, neither fair nor unfair competition, somewhat unfair competition, and entirely unfair competition. In the analysis, we code this question with values 0–4, where higher values mean more unfair competition.
- *Perceived negative image*: this question asked: “Overall, how do you feel about [company] as a company?” using a *five-point favorability scale*: very favorable, somewhat

¹³Table B.1 provides the market definitions used for all five cases.

¹⁴In two of the cases, we adjusted the language to fit the company’s context more closely: for Apple, the statement refers to “higher prices or reduced innovation”; for Meta, it refers to “decreased data security or increased ads.”

favorable, neutral, somewhat unfavorable, and very unfavorable.¹⁵ In the analysis, we code this question with values 0–4, where higher values mean more negative image.

2.5 Information Treatments

Within each treatment arm, the information followed a common structure and type of argument across cases, while the specific facts and examples were tailored to the assigned case. Appendix E reproduces the screenshots for each information treatment and case. The treatment screens were fact-based and included links to the underlying sources to make the information credible and easy to verify. Across treatments, the sources included market-share estimates, newspaper articles, academic research, expert surveys, reports, litigation summaries, and short excerpts from news programs or documentary-style segments, such as CBS Morning News or 60 Minutes, summarizing allegations about the assigned case.

We summarize the five information conditions below. Each respondent was randomly assigned to one of the following groups, with equal probability:

- *Control group*: respondents did not receive any additional information, only a short screen stating that they had been randomly selected not to receive information.
- *Market-share treatment*: respondents received additional information about the assigned company’s share of the relevant U.S. market. The market definition matched the one used in the prior-belief question. The information stated that Google has an estimated 84 percent share of U.S. online search advertising, Live Nation controls roughly 80 percent or more of primary ticketing at major U.S. concert venues, Apple has an estimated 84 percent share of the U.S. mobile application market, Luxottica controls at least 60 percent of the U.S. eyewear market, and Meta has a 72 percent share of the U.S. social media market.
- *Consumer-harm treatment*: respondents received additional information arguing that the company’s market power may harm consumers. Depending on the case, the treatment described channels such as higher prices, higher fees, reduced quality, reduced innovation, reduced data security, increased ads, or fewer choices. For example, in the Google case, the treatment explained that although Google Search appears free to users, advertisers pay for search ads and may pass those higher advertising costs on to consumers through the prices of products and services. The treatment also linked to a

¹⁵Following Perez-Truglia and Yusof (2026), we also included an exploratory product-choice task designed to measure company image in a more revealed-preference way. Respondents chose between a backpack with the assigned company or brand logo and an otherwise similar backpack without that logo in several price scenarios. See Appendix B.7 for more details.

report discussing Google’s market power and cited Clark Center survey evidence from leading economists about whether Google’s dominance harms consumers and society.

- *Unfair-competition treatment:* respondents received information suggesting that the assigned company’s dominance may reflect conduct or advantages that could be perceived as unfair. While fairness judgments can be subjective, we do not define the conduct as fair or unfair ourselves; rather, each treatment focused on conduct that media coverage or other parties had characterized as unfair. These treatments focused on practices that may make it harder for rivals to compete, such as default arrangements, exclusive contracts, restrictions within a platform ecosystem, vertical control over distribution, or acquisitions of potential competitors. For example, in the Live Nation case, the treatment discussed an academic article by Michael Carrier, a law professor at Rutgers University, arguing that Live Nation and Ticketmaster used exclusive ticketing contracts and pressure on venues to reinforce Ticketmaster’s position. The treatment described allegations that venues could lose access to Live Nation concerts if they used rival ticketing services, and that multiyear exclusive contracts made it difficult for competing ticketing firms to reach major events.
- *Negative-image treatment:* respondents received information about a recent scandal or controversy involving the assigned company that was not itself an antitrust allegation. These treatments were designed to affect respondents’ broader perception of the company rather than to provide direct information about market power or consumer harm. For example, in the Apple case, the treatment described a Wall Street Journal article about a proposed class-action lawsuit alleging gender-based pay discrimination. The treatment explained that the plaintiffs sought to represent thousands of women employed by Apple and argued that the company’s hiring and performance-review practices contributed to persistent pay disparities.

To reduce the possibility that respondents would make unintended inferences from receiving information, the randomization was made explicit: respondents were told that some participants would be randomly selected to receive information and that they would learn on the next screen whether they had been selected.¹⁶

¹⁶For example, a respondent receiving information about market share could wrongly infer that they were receiving the information because their prior belief about market share was very inaccurate. By making the randomization explicit, we mitigate this concern.

2.6 Posterior Beliefs

Immediately after the information-provision screen, regardless of whether the respondent was assigned to the control group or one of the four treatment groups, we re-elicited the same four belief questions from the prior-beliefs module. For example, respondents assigned to Google were again asked the four Google-specific belief questions described above. Comparing posterior beliefs to prior beliefs allows us to measure belief updating about the assigned company and case.

As discussed above, the four beliefs are empirically interdependent: each information treatment may affect beliefs beyond the one it was designed to target. For example, information that a company used unfair practices to gain market power could affect not only perceived unfairness, but also beliefs about the company’s market share, consumer harm, and negative image. Measuring all four belief updates is therefore central to the empirical strategy, which uses these updates to study the relative contribution of each belief to demand for antitrust enforcement.

We also took steps to reduce the possibility that respondents would make unintended inferences from the repeated belief questions. Before eliciting posterior beliefs, we informed respondents that all participants were asked the same questions again, regardless of their earlier responses or whether they received any information. This wording was intended to reduce the chance that respondents would interpret the repeated belief questions as a signal that their initial answers were incorrect.

2.7 Outcomes

To understand how information changes demand for antitrust enforcement, we included several outcome questions. Table B.2 reports their wording and scales, which are summarized below. The main outcomes are case-specific measures of support for the plaintiff side and for remedies in the assigned case. The variable *plaintiff support* captures which side respondents support in the assigned case. For example, respondents assigned to the Google case were asked: “Imagine you were appointed as the judge in charge of deciding in the case between Google and the Department of Justice (DOJ). Which side would you be more likely to support?”¹⁷ The five response options were strongly support the company, somewhat support the company, support neither side, somewhat support the plaintiff side, and strongly support the plaintiff side. We refer to this as the *five-point side-taking scale*. In the analysis, we code this outcome on a 0–4 scale, where higher values mean stronger support for the plaintiff side

¹⁷The wording was tailored to each assigned case. For example, the counterpart could be the Federal Trade Commission, Epic Games, or a class action lawsuit rather than the Department of Justice.

in the assigned case. When a question asks respondents whether they support or oppose a remedy or policy proposal, we use a *five-point support scale*: strongly oppose, somewhat oppose, neither support nor oppose, somewhat support, and strongly support. The remedy outcomes measure the *intensity* of demand for antitrust enforcement: whether respondents support more consequential structural or conduct remedies, such as breaking up the company, rather than only a favorable ruling against it (Kwoka and Moss, 2012).¹⁸ The first remedy, *break-up remedy*, is a break-up or divestiture remedy. For example, respondents assigned to Live Nation were asked whether they supported a measure that would “Require Live Nation to sell off Ticketmaster so it becomes a separate company.” The second remedy, *conduct remedy*, would restrict a specific business practice. In the Live Nation case, for example, respondents were asked whether they supported a measure that would “Prevent Live Nation from making exclusive ticketing deals with venues and artists that require them to use Ticketmaster.” In the analysis, we code each remedy outcome on a 0–4 scale, where higher values indicate stronger support for the remedy. For conciseness, we combine the two remedy outcomes into a single *remedy index*, defined as their average and therefore also measured on the 0–4 scale.

Beyond measuring demand for antitrust in the assigned case, the survey was designed to measure the *breadth* of the treatment effects: whether respondents extrapolate from a specific antitrust case to beliefs and preferences about antitrust more generally. For example, if information about Google increases the belief that Google harms consumers, we ask whether it also increases the belief that market power by dominant companies more broadly harms consumers. Likewise, if information about Google increases demand for antitrust enforcement in the Google case, we ask whether it spills over to support for antitrust enforcement and competition policy in other concentrated markets.

For brevity, when we discuss prior beliefs or posterior beliefs without further qualification, we refer to beliefs about the assigned company and case. When discussing beliefs about dominant companies more broadly, we explicitly call them *general* beliefs. The survey introduced the general-beliefs module by asking respondents to think about “companies that have a dominant position in their markets”. To make this idea concrete without repeating the assigned case, the survey used the other four cases as examples of dominant companies.¹⁹ We then elicited a set of *general posterior beliefs* that paralleled the four belief questions about the assigned case. Respondents were asked the average share of the relevant U.S.

¹⁸To mitigate order effects, the survey randomized the order in which respondents saw the two remedy questions.

¹⁹For example, a respondent assigned to the Google case saw examples of Apple in smartphones and mobile operating systems, Meta in social-media advertising, Luxottica in eyeglass frames and lenses, and Live Nation in ticketing and live-event promotion.

markets held by these dominant companies, whether they agreed that the market power of dominant companies harms consumers, whether the market power of these dominant companies is mainly the result of fair or unfair competition, and how they felt overall about these dominant companies.

Because remedies are context-specific, the broader analogue for the general outcomes is support for antitrust enforcement and competition policy. We take multiple approaches to elicit this broader form of support. First, the *general antitrust support* outcome asks whether antitrust laws should be enforced more leniently, more strictly, or about the same as they are now. We refer to this as the *five-point enforcement-strictness scale*, coded from much more leniently (0) to much more strictly (4). Second, we described three real-world policy proposals related to antitrust and competition policy and asked respondents whether they supported or opposed each proposal using the five-point support scale.²⁰ The *support EO* question described Executive Order 14036, “Promoting Competition in the American Economy”, which directs federal agencies to promote competition and strengthen antitrust enforcement. The *support conduct bill* question described the American Innovation and Choice Online Act, which would restrict certain business practices by large online platforms that may disadvantage competing businesses. The *support merger bill* question described the Platform Competition and Opportunity Act, which would make it harder for very large online platforms to acquire other companies. Before these policy questions, respondents were informed that, at the conclusion of the study, we would share the anonymous survey results with politicians and organizations relevant to these issues. This type of consequentiality script has been shown to increase respondents’ perceptions that their answers may matter for real-world decisions (e.g., Elías et al., 2019).

We also include three behavioral measures of policy support. For *petition signing*, respondents read about a petition calling for stronger antitrust and competition enforcement and indicated whether they would be interested in signing it. Respondents who answered yes were shown a link to the petition website. In addition to asking about their willingness to sign the petition, we recorded whether they clicked the petition link. For *contact representative*, respondents indicated whether they would consider contacting an elected official to express concerns about abuses of market power by companies such as the one assigned to them. Those who answered yes were shown an editable draft message advocating fair competition and stronger antitrust enforcement. Based on each respondent’s state of residence, the survey identified one of their U.S. senators and provided a link to the senator’s contact page. We recorded whether respondents copied the suggested message and whether

²⁰The order of the three policy-proposal questions was randomized.

they clicked through to the contact page.²¹ Finally, for *antitrust donation*, respondents allocated a \$100 donation between the American Antitrust Institute, a nonprofit organization that promotes competition through research, education, and advocacy, and World Relief, an international humanitarian organization. Respondents could allocate any amount from \$0 to \$100 to either organization, provided that the allocations summed to \$100. To make the decision consequential, one survey participant was randomly selected, and the \$100 was donated according to that participant’s allocation. We measure the amount allocated to the American Antitrust Institute.

Because the six policy-support metrics—the three policy-proposal items, petition signing, contact representative, and antitrust donation—are elicited on different scales, we combine them into a *policy index* following the approach of Kling et al. (2007): we orient each component so that higher values indicate greater support for antitrust policy, standardize each component using the control-group mean and standard deviation, average the standardized components, and then standardize the average again using the control-group mean and standard deviation.²²

2.8 Follow-Up Survey

One month after the baseline survey, respondents were invited to complete another survey, which we refer to as the follow-up survey. The main goal of this survey was to measure the *durability* of treatment effects: whether any changes in beliefs and demand remained visible about one month later.

The follow-up survey instrument is attached as Appendix G and described in detail below. Panel B of Figure 1 summarizes the structure of this survey. We designed the beginning of the follow-up survey to reduce demand effects. Before returning to the antitrust questions, respondents answered a short set of filler questions about their use of and attitudes toward artificial intelligence. The follow-up survey also used University of Zurich branding rather than the UCLA branding from the baseline survey. Next, the survey included a brief summary of the same case that was assigned at baseline. It then re-elicited the same belief measures and most of the same outcome measures from the baseline survey, except for the behavioral policy-support outcomes: petition signing, contacting a representative, and the donation-allocation task. We excluded these outcomes for two reasons. First, we wanted to keep the follow-up survey short enough to accommodate the trial-simulation module described below.

²¹Clicking the petition link, copying the suggested message, or visiting the senator’s contact page does not establish that a respondent ultimately signed the petition or submitted the email. We therefore use these follow-through measures as behavioral validation of the stated-choice outcomes.

²²This standardization is not needed for the remedy index, which averages two outcomes elicited using the same five-point support scale.

Second, it would be somewhat odd to re-elicited the same behavioral outcomes—for example, respondents who had already signed the petition at baseline would not be eligible to sign it once again.²³

At the end of the follow-up survey we included a real-time trial simulation. In the Google case, for example, the respondent was asked to take the role of a Department of Justice lawyer, while an AI lawyer represented Google. Respondents exchanged a few messages with the AI lawyer in real time, and an AI judge evaluated which side made the more convincing case. To encourage respondents to take the exercise seriously and put effort into their arguments, respondents were told that 10% of participants would be randomly selected for an additional bonus based on how persuasive their arguments were. By placing this simulation at the end of the survey, we ensured that the information exchanged with the AI lawyer would not contaminate the rest of the outcomes.²⁴ Because the trial simulation generated unstructured text data, analyzing these responses requires additional time. We therefore leave the analysis of this module to a future version of the paper.

2.9 Complementary Surveys

We conducted two supplemental surveys with different samples of respondents. The first one is referred to as the “expert forecast survey.” The goal of this survey is to assess the extent to which the experimental results were predictable or surprising. A sample of academic and professional experts were provided with a detailed description of the main experiment, including the antitrust cases, information treatments, and main outcomes. They were then asked to forecast the average treatment effects and to rank the potential causal mechanisms. Appendix D provides more details on the expert sample and survey design, and Appendix I reproduces the full survey instrument. The second supplemental survey is referred to as the “emotions survey.” The goal of this survey is to evaluate respondents’ emotional reactions to the information treatments. Appendix C provides more details on the survey design and implementation, and Appendix H reproduces the full survey instrument.

²³Likewise, if respondents had already contacted a representative at baseline, it would be odd to ask them again to contact a representative.

²⁴For example, the conversation with the AI could expose respondents to arguments or information related to a different treatment as the one they were originally assigned to.

3 Sample and Baseline Demand for Antitrust

3.1 Implementation Details

The experiment was pre-registered at the AEA RCT Registry ([AEARCTR-0018249](#)) and was reviewed and approved by the North General Institutional Review Board at UCLA. We conducted the experiment with a sample of American respondents recruited through Prolific. We required respondents to be U.S. residents and to complete the survey on a laptop or desktop computer. The baseline survey was fielded on May 4–5, 2026. A total of 4,016 respondents completed the baseline survey. The median completion time was 16.7 minutes and respondents were offered \$3 for completion. We invited the baseline respondents to participate in the follow-up survey approximately one month later. A total of 2,865 respondents completed the follow-up survey, corresponding to 71.3% of the baseline sample.²⁵ The median completion time was 14.5 minutes.²⁶

Given growing concerns about online data quality and AI-assisted survey responses, we implemented several quality-control measures. The baseline survey informed respondents that the use of AI tools or automated assistance was prohibited. Prolific has its own measures in place to detect and prevent AI use and automated AI agents in surveys ([Prolific, 2026](#)). According to independent accounts, at the time that we fielded the survey, the estimated share of Prolific participants using AI was reported to be very small (e.g., [Çelebi et al., 2026](#); [Affonso, 2026](#)). However, out of an abundance of caution, we included our own independent AI detection tools. The main AI-use detection, developed by [Çelebi et al. \(2026\)](#), uses a short video that AI agents can easily miss because they ingest snapshots of the screen rather than real-time video feed. Indeed, this AI detection tool was found to be highly effective at detecting AI agents ([Çelebi et al., 2026](#)). This video-code check was passed by 98.5% of respondents. The 1.5% who failed provide an upper bound for the potential share of responses generated by AI agents, since human respondents could fail this check if they were not paying close attention. Thus, this AI check confirms that the vast majority of the respondents were not using AI tools to respond to our survey. The rest of the data quality checks reinforce this view. We also included CAPTCHA verification, and the mean reCAPTCHA score was 0.96 on a 0–1 scale, where higher values indicate a greater likelihood that the respondent is human.²⁷ Finally, an open-ended question in the baseline survey allowed us to track

²⁵Because respondents could take a few days to respond to the follow-up invitation, the elapsed time between the baseline and follow-up surveys varies somewhat across individuals. On average, the follow-up survey was completed 29.9 days after the baseline survey.

²⁶Respondents were offered a minimum of \$2 for completing the follow-up survey.

²⁷We included both a reCAPTCHA verification challenge and the video-check in the follow-up survey too, with similar results.

respondents’ typing behavior. In total, 92.5% of respondents entered text; among them, 99.5% typed at speeds consistent with human input. These patterns again suggest that the share of AI-assisted or AI-agent-generated responses was very low.

We included standard attention checks, which 96.3% of respondents passed even though the checks were placed at the end of the survey, when survey fatigue was likely highest. Among baseline respondents, 98.3% rated the survey questions as easy or very easy. The distribution of completion times suggests little concern about substantial speed-taking: only 4.5% of respondents completed the survey in under eight minutes.

The survey also collected several background measures to characterize respondents and assess their familiarity with the assigned companies and antitrust cases. Respondents entered the survey with different levels of familiarity with the assigned companies and cases. Most respondents had heard of Google, Apple, Meta, and Live Nation, while Luxottica was much less familiar. Awareness of the specific antitrust disputes was lower and varied more across cases, from 5.8% for Luxottica to 43.9% for Live Nation; Appendix B.3 provides more detail.²⁸ We also asked respondents whether they perceived the survey as politically biased; 86.7% reported no bias, 10.6% a left-leaning bias, and 2.8% a right-leaning bias.

3.2 Sample Characteristics and Randomization Balance

Table 1 reports descriptive statistics and randomization balance. Column (1) describes the full baseline sample: 58% of respondents are female, 46% are 35 or younger, 66% are White, 19% are Black, 6% are Asian, and 6% are Hispanic. About 66% of respondents report annual household income above \$50,000, 55% have a college degree, 43% identify as Democrats, and 25% identify as Republicans. As is common in online samples, our sample differs somewhat from the general U.S. population: respondents are younger, more educated, more female, and more Democratic-leaning than the benchmark population—see Appendix B.2 for more details. Columns (2)–(6) report the same characteristics separately by information-treatment arm: control, market share, consumer harm, unfair competition, and negative image. Column (7) reports the p-value from a joint test of equality across treatment arms. Observable characteristics are balanced across treatment groups: the differences are small in magnitude, and none of the reported p-values indicates statistically significant imbalance. The final panel reports follow-up participation. Overall, 71% of baseline respondents completed the follow-up survey, and follow-up participation is similar across treatment arms, suggesting no meaningful selective attrition by information treatment.

²⁸Prior awareness of the antitrust cases, if anything, biases the results against detecting treatment effects because respondents who were already familiar with a case received less new information from the information treatment.

3.3 Baseline Beliefs

We begin by describing baseline beliefs, using respondents assigned to the control group. The top row of Figure 2 summarizes these results. Panel A shows that the average perceived market share was 67.2%. This average masks substantial heterogeneity: the 10th percentile is 40.0%, while the 90th percentile is 86.9%. Panel B shows that 68.3% of respondents somewhat or strongly agree that the assigned firm’s market power harms consumers; again, these averages mask meaningful heterogeneity, as these beliefs are stronger for some subjects than for others. Panel C shows that 53.8% of respondents view the firm’s market power as arising more from unfair than fair competition. And Panel D shows that 38.8% report an unfavorable view of the assigned company. The corresponding general beliefs about dominant firms show similar patterns.²⁹ The four beliefs are correlated with one another, though far from perfectly so. In the control group, pairwise correlations range from 0.01 between perceived market share and perceived negative image to 0.71 between perceived unfair competition and perceived negative image, with an average correlation of 0.37 across the six belief pairs. This pattern is natural given that the beliefs are interdependent: for example, respondents who perceive a higher market share also tend to perceive somewhat greater consumer harm.

3.4 Baseline Demand for Antitrust Enforcement

We next describe the demand for antitrust enforcement in the control group, to provide a sense of the baseline support. The main pattern is that baseline demand is substantial. The bottom row of Figure 2 reports selected outcomes for the control group. Panel E shows that baseline plaintiff support is substantial but far from unanimous: 56.4% of control-group respondents somewhat or strongly support the plaintiff side, with an average response of 2.46 on the five-point side-taking scale. Support for remedies depends on the type of remedy. Although many respondents favor antitrust enforcement, support is lower for the more punitive structural remedy than for restrictions on specific business practices: while 46.9% somewhat or strongly support the break-up remedy (Panel F), 70.1% support the conduct remedy (Panel G).

Support is also high when respondents are asked about antitrust enforcement and competition policy more generally. Panel H shows that 75.5% of control-group respondents somewhat or strongly support increasing enforcement of antitrust laws in concentrated industries. Support for the three policy proposals is high, with more than 80% of control-group respondents somewhat or strongly supporting each proposal. By contrast, the behavioral measures

²⁹Respondents perceive dominant firms as having large market shares and tend to agree that market power harms consumers, while views about unfair competition and negative image are more mixed—see Appendix B.4 for details.

show more moderate support. On average, respondents allocate 39.6% of the hypothetical donation budget to the antitrust organization. And although 49.5% say they would sign the petition and 40.6% say they would consider contacting their representative, the behavioral proxies, such as clicking the relevant links, indicate that smaller shares took subsequent steps toward these actions; see Appendix B.6 for details. Additional baseline outcomes are discussed in Appendix B.4. Although baseline beliefs and support vary somewhat across cases, the overall pattern is consistent: respondents perceive the assigned firms as having substantial market power and express meaningful support for the plaintiff side, case-specific remedies, and general antitrust enforcement; see Appendix B.8 for more details.

Although we observe a high level of support for antitrust enforcement on average, baseline beliefs and preferences vary considerably across respondents. Before interpreting this cross-sectional variation as meaningful heterogeneity, we assess the extent to which it may reflect measurement error. We use two test-retest exercises among respondents assigned to the control group. First, we compare prior and posterior beliefs elicited within the baseline survey. Because control-group respondents received no information between these elicitation, the resulting correlations provide a measure of short-run reliability. The correlations are high, ranging from 0.88 to 0.95 across the four belief measures. Under the classical measurement-error model, these correlations can be interpreted as the share of observed variation attributable to signal rather than noise. The high correlations therefore indicate that variation in the belief measures contains a strong and meaningful signal.

However, these within-survey correlations may overstate reliability if some respondents remembered and repeated their earlier answers. To address this concern, we conduct a second test by comparing responses across the baseline and follow-up surveys. This longer-run comparison is more demanding because respondents may have consumed additional information about the companies or antitrust cases and genuinely revised their views between survey waves. Thus, the cross-wave correlations reflect both measurement error and genuine changes in beliefs and preferences. The test-retest correlations range from 0.57 to 0.78 for the belief measures and from 0.60 to 0.64 for the three main case-specific outcomes. These correlations suggest that the survey measures capture relatively stable views rather than pure noise, especially given that the two surveys were conducted one month apart. They are also in the range of test-retest correlations reported for other subjective survey measures, including life satisfaction and economic preferences (Krueger and Schkade, 2008; Dohmen and Jagelka, 2024)—for more details, see Appendix B.12.

Finally, we examine how prior beliefs relate to demand for antitrust enforcement. Each of the four prior beliefs is positively and significantly correlated with baseline demand for antitrust enforcement—see Appendix B.5 for more details. Plaintiff support is strongly as-

sociated with the beliefs that the firm’s market power harms consumers, reflects unfair competition, and gives respondents a more negative view of the company. The correlation with perceived market share is positive but more modest. These correlations, of course, do not imply causation, which is why we use an experimental approach.

4 Average Treatment Effects

4.1 Treatment Effects on Beliefs

Figure 3 reports treatment effects on the four belief measures. As a placebo check, Appendix B.13 shows that prior beliefs are balanced across treatment arms before respondents learned whether they would receive information. Let $K = \{MS, CH, UC, NI\}$ denote perceived market share, perceived consumer harm, perceived unfair competition, and perceived negative image. For respondent i and belief $k \in K$, let b_{ik}^{prior} and b_{ik}^{post} denote prior and posterior beliefs. Define the belief update as $\Delta b_{ik} \equiv b_{ik}^{post} - b_{ik}^{prior}$. Let T_i^k indicate assignment to information treatment k , with $T_i^k = 0$ for all k in the control group. Let s_{ic}^{MS} denote the market-share signal corresponding to respondent i ’s assigned case. We start with the market-share treatment because it provides the cleanest illustration of belief updating: unlike the other treatments, it gives respondents a numeric signal on the same scale as the belief elicitation. Panel A shows that the market-share treatment increases perceived market share by 6.8 pp. This effect is highly statistically significant and large in magnitude. The positive average effect is also consistent with the distribution of prior beliefs: among respondents assigned to the market-share treatment, prior beliefs were, on average, 9.8 pp below the corresponding signal.

We use a simple framework of Bayesian learning as a lens to interpret these effects on beliefs (e.g., Cavallo et al., 2017; Cullen and Perez-Truglia, 2022). In the standard case in which the respondent’s prior belief and the signal are normally distributed with known variances, the posterior mean is a precision-weighted average of the prior mean and the signal (e.g., DeGroot, 1970). Subtracting the prior from both sides implies that belief updating is proportional to the gap between the signal and the respondent’s prior belief:

$$b_{i,MS}^{post} - b_{i,MS}^{prior} = \alpha_{MS} (s_{ic}^{MS} - b_{i,MS}^{prior}) + \varepsilon_{i,MS}. \quad (1)$$

The parameter α_{MS} measures the degree of updating. If $\alpha_{MS} = 0$, respondents ignore the information; if $\alpha_{MS} = 1$, respondents fully adjust to the signal. Equation (1) implies that respondents whose priors are below the signal should revise their beliefs upward, while respondents whose priors are above the signal should revise their beliefs downward.

In practice, respondents may also revise their answers between the prior and posterior elicitations even when they do not receive the market-share treatment. For example, they may reconsider the question, correct an earlier typo, or move their answer in the direction of the true value for spurious reasons unrelated to the treatment. Following the standard approach in information-provision experiments, we separate treatment-induced learning from this spurious updating by comparing respondents who were shown the information to those who were not:

$$\Delta b_{i,MS} = \tau_{MS} + \delta_{MS} T_i^{MS} + \alpha_{MS} T_i^{MS} (s_{ic}^{MS} - b_{i,MS}^{prior}) + \eta_{MS} (s_{ic}^{MS} - b_{i,MS}^{prior}) + \varepsilon_{i,MS}. \quad (2)$$

The coefficient η_{MS} captures spurious updating: the relationship between the market-share gap and belief revisions among respondents in the control group. The coefficient α_{MS} captures how the market-share treatment changes the slope of belief updating with respect to the prior gap, above and beyond that spurious updating.

Panel B of Figure 3 displays the empirical analogue of equation (2). The y-axis is belief updating, $\Delta b_{i,MS}$, and the x-axis is the prior gap, $s_{ic}^{MS} - b_{i,MS}^{prior}$, plotted separately for the market-share treatment and control groups. The control-group slope corresponds to η_{MS} , while the treatment-group slope corresponds to $\alpha_{MS} + \eta_{MS}$; the difference between the two slopes is therefore α_{MS} . At any given prior gap, the treatment-induced belief update is $\delta_{MS} + \alpha_{MS} (s_{ic}^{MS} - b_{i,MS}^{prior})$. Consistent with the learning framework, respondents update upward when their priors are below the signal and downward when their priors are above it. Empirically, the treatment and control lines cross close to zero, indicating little treatment-induced updating among respondents whose priors were already close to the signal.

Next, we turn to the consumer-harm treatment. Panel C of Figure 3 shows that this treatment increased perceived consumer harm by 0.52 points on the 0–4 scale, an effect that is both statistically significant ($p < 0.001$) and economically meaningful. Panel D shows the heterogeneity in belief updating by prior beliefs. The consumer-harm treatment differs from the market-share treatment because the signal was implicit rather than explicit. That is, the consumer-harm treatment did not tell respondents that the firm’s consumer harm was equal to a specific number on the 0–4 scale. Instead, it was up to the respondents to interpret the information as a value of 3, 3.5, 4, or some other value. The fact that the average update was positive implies that, on average, the implied signal was above the average prior belief (2.7 on the 0–4 scale).

To analyze belief updating when signals are implicit rather than explicit, we estimate a version of equation (2) that replaces the observed prior gap with the prior belief itself. The

estimable regression for $k = CH$ is:

$$\Delta b_{i,CH} = \tau_{CH} + \delta_{CH} T_i^{CH} + \alpha_{CH} T_i^{CH} b_{i,CH}^{prior} + \eta_{CH} b_{i,CH}^{prior} + \varepsilon_{i,CH}. \quad (3)$$

The control-group slope is η_{CH} , while the treatment-group slope is $\eta_{CH} + \alpha_{CH}$; the difference between the two slopes is therefore α_{CH} . At any given prior belief, the treatment-induced belief update is $\delta_{CH} + \alpha_{CH} b_{i,CH}^{prior}$. If the treatment moves respondents toward an implied signal, the treatment and control lines should cross at the value of the prior belief equal to that implied signal. Panel D of Figure 3 shows that the treatment and control lines cross close to 4. This suggests that respondents interpreted the consumer-harm narrative as a strong signal near the top of the scale: respondents with priors below 4 update upward, while respondents already near 4 have no reason to update downward.

Panels E and F of Figure 3 repeat the exercise for perceived unfair competition. Panel E shows that the unfair-competition treatment shifted perceived unfair competition upward, on average by 0.60 points on the 0–4 scale. Since this treatment does not make the signal explicit, we estimate the analogue of equation (3) with $k = UC$. Panel F shows that respondents with lower prior unfairness beliefs update upward by more, while the treatment-control gap narrows as prior beliefs move toward the top of the scale. The fitted lines converge near the top of the scale, suggesting that respondents interpreted the unfair-competition treatment as a strong signal near the top of the unfairness scale.

Panels G and H turn to perceived negative image. Panel G shows that the negative-image treatment increases the perceived negative image of the assigned company by an average of 0.51 points on the 0–4 scale. In Panel H, we estimate the implicit-signal specification with $k = NI$. The treatment and control groups are nearly identical among respondents with prior beliefs equal to 4. This suggests that respondents interpreted the negative-image treatment as a strong signal near the top of the negative-image scale: respondents with lower priors update upward, while respondents already near the top have no reason to update downward.

Table 2 reports the average treatment effects on the main beliefs and outcomes. Panel A reports raw coefficients in the original units of each outcome. To make it easier to compare effect sizes across outcomes, we focus mostly on Panel B, which reports coefficients standardized by the standard deviation of the corresponding outcome. Columns (1)–(4) report the effects on all four case-specific belief measures. The diagonal coefficients correspond to the effects discussed in Figure 3: each treatment significantly shifts the belief it was designed to target. The off-diagonal panels shows that, additionally, the treatments often spills over to other beliefs.

The market-share treatment is the most compartmentalized: it increases perceived mar-

ket share by 0.39 SD, but its only statistically significant spillover is on perceived unfair competition, at 0.13 SD. Our preferred interpretation is that, after learning that the firm has a larger market share than they thought, respondents infer that the company may have done something unfair to gain that competitive advantage.

For the other treatments, the spillovers are broader. The consumer-harm treatment increases perceived consumer harm by 0.48 SD, but it also increases perceived market share by 0.34 SD, perceived unfair competition by 0.41 SD, and perceived negative image by 0.44 SD. This suggests that information about consumer harm changes respondents’ broader interpretation of the case: if the firm’s market power harms consumers, respondents also infer greater dominance, less fair conduct, and a more negative overall view of the company.

The unfair-competition treatment has similarly broad effects. It increases perceived unfair competition by 0.53 SD, while also increasing perceived market share by 0.36 SD, perceived consumer harm by 0.38 SD, and perceived negative image by 0.44 SD. This pattern is consistent with respondents interpreting unfair conduct as evidence of broader market power and consumer harm: if the firm uses unfair practices to preserve its position, respondents infer that it is also more dominant, more harmful to consumers, and deserving of a more negative evaluation.

The negative-image treatment primarily affects perceived negative image, increasing it by 0.42 SD. It also has significant spillovers to perceived consumer harm (0.19 SD) and perceived unfair competition (0.21 SD), though no detectable effect on perceived market share. This suggests that information portraying the firm’s leadership as unethical leads respondents to draw broader inferences about the firm’s conduct: they become more likely to believe that the company harms consumers and competes unfairly, but not necessarily that it has a larger market share.

These patterns reinforce the idea that the four beliefs are interdependent: information designed to move one belief can also change how respondents think about market power, consumer harm, unfair conduct, and the company more generally. In this section, we focus on the average treatment effects; in Section 5, we return to this belief updating and use it to disentangle the causal effects of each belief.

4.2 Treatment Effects on Demand for Antitrust Enforcement

Column (5) of Table 2 reports effects on plaintiff support, our most direct measure of which side respondents favor in the assigned case. Support for the plaintiff side responds most strongly to information about consumer harm and unfair competition, while market share alone has little effect. Despite the strong effect on the perceived market share, the market-share treatment has small (0.09 SD) and only marginally significant ($p=0.066$) effect on

the plaintiff support. By contrast, the other three treatments increase plaintiff support substantially. The consumer-harm and unfair-competition treatments produce the largest effects, at 0.45 SD and 0.48 SD, respectively (both $p < 0.001$). The negative-image treatment has a somewhat smaller effect on plaintiff support, at 0.29 SD ($p < 0.001$).

We organize the remaining outcome effects around three dimensions of demand for antitrust enforcement: intensity, breadth, and durability. Intensity asks whether information moves respondents beyond plaintiff support toward support for concrete remedies, such as breaking up the company. Breadth asks whether information about the assigned case spills over to beliefs and preferences about dominant firms more broadly. Durability asks whether the effects remain visible in the follow-up survey about one month later.

Column (6) of Table 2 shows the effects on intensity. The consumer harm and unfair competition treatments produce the strongest evidence that respondents are willing to move beyond agreement with the plaintiff side toward more consequential remedies. The market-share treatment has no statistically significant effect on the remedy index (0.06 SD, $p = 0.262$). By contrast, the consumer-harm treatment increases the remedy index by 0.28 SD, while the unfair-competition treatment increases it by 0.33 SD (both $p < 0.001$). The negative-image treatment also increases remedy support, but the effect is smaller, at 0.19 SD ($p < 0.001$). Most importantly, this is just the average effect of the negative image treatment; Section 5 asks whether it operates through perceived negative image itself or through the treatment’s spillovers to perceived consumer harm and unfair competition.

Columns (1)–(6) of Table 3 speak to breadth. Overall, breadth is strongest for consumer harm, intermediate for unfair competition, and limited for negative image and market share. The consumer-harm treatment increases all four general beliefs, with effects of 0.13 SD on general market share ($p = 0.009$), 0.25 SD on general consumer harm ($p < 0.001$), 0.17 SD on general unfair competition ($p < 0.001$), and 0.15 SD on general negative image ($p = 0.002$). It also increases general antitrust support by 0.13 SD ($p = 0.010$). The effect on the policy index is smaller, at 0.05 SD, and not statistically significant ($p = 0.270$). The 90% confidence interval, $[-0.03, 0.13]$, does not rule out modest positive effects. Still, the smaller point estimate suggests that consumer-harm information may move general support for stricter enforcement more than support for specific policy instruments or behavioral actions, perhaps because some respondents recognize the problem but are less enthusiastic about the available political or policy solutions (Kuziemko et al., 2015; Stantcheva, 2021). The unfair-competition treatment also has some breadth, increasing general consumer harm by 0.16 SD ($p = 0.001$) and general unfair competition by 0.12 SD ($p = 0.017$), but its effects on general antitrust support and the policy index are not statistically significant (0.08 SD, $p = 0.110$, and -0.05 SD, $p = 0.317$). Our preferred interpretation is that learning that executives at one com-

pany engaged in unfair practices may not lead respondents to infer that executives at other dominant firms behave similarly. The negative-image treatment has more limited breadth: it has small, marginally significant effects on general consumer harm (0.09 SD, $p=0.071$) and general unfair competition (0.09 SD, $p=0.077$), but not on general negative image (0.06 SD, $p=0.204$), general antitrust support (0.08 SD, $p=0.110$), or the policy index (-0.01 SD, $p=0.807$). In other words, respondents appear willing to update their view of the specific company, but do not generalize from misconduct at one firm to a broader negative view of dominant firms as a class. The market-share treatment again mostly affects beliefs about market share itself, increasing general perceived market share by 0.16 SD ($p=0.001$), with no corresponding increase in general antitrust support (-0.00 SD, $p=0.960$) or the policy index (-0.04 SD, $p=0.475$). This reinforces the main interpretation from the case-specific outcomes: respondents update their beliefs about dominance, but absent corresponding concerns about consumer harm or unfair conduct, they do not treat dominance itself as a sufficient reason to demand stronger antitrust enforcement.

Columns (11) and (12) of Table 2 speak to durability. The consumer-harm treatment remains the most durable: one month later, it still increases plaintiff support by 0.15 SD ($p=0.010$) and the remedy index by 0.12 SD ($p=0.032$). The unfair-competition treatment also has persistent effects, but they are somewhat smaller and less precisely estimated, with effects of 0.11 SD on plaintiff support ($p=0.054$) and 0.10 SD on the remedy index ($p=0.082$). By contrast, the negative-image treatment has little durability: its follow-up effects are 0.06 SD for plaintiff support ($p=0.277$) and 0.01 SD for the remedy index ($p=0.839$). The market-share treatment continues to have no meaningful effect on the main outcomes, with follow-up effects of 0.01 SD and 0.02 SD ($p=0.904$ and $p=0.700$).

One interpretation is that durability declines with inferential distance from the treatment: information about a specific firm may remain relevant for views about that firm, while broader beliefs and policy preferences require an additional extrapolation that fades more quickly. Columns (7)–(12) of Table 3 examine whether the broader effects persist in the follow-up survey. Unlike in the baseline survey, the follow-up effects on general beliefs and general antitrust support are not statistically significant. These null effects should be interpreted cautiously: the follow-up effects are substantially attenuated, and the follow-up sample is 29% smaller than the baseline sample. As a result, the confidence intervals are too wide to rule out modest effects that are both broad and durable. The one unusual pattern is that column (12) shows statistically significant negative effects on the follow-up policy index for all four treatments. We do not interpret these estimates as evidence that the treatments reduced support for antitrust policy, since they are not accompanied by negative effects on general beliefs or general antitrust support. A more plausible explanation is that, by chance,

the control group scored higher on the follow-up policy-index components.

More broadly, the follow-up estimates line up directionally with the baseline estimates, but they are smaller in magnitude. This pattern is consistent with other information-provision experiments, where the effects of information tend to decay over time, probably because, in the interim, some subjects forget the experimental information or receive new information. For example, [Cavallo et al. \(2017\)](#) find that 45.6% of the belief updating induced by information about inflation persisted four months later. In our setting, a pooled comparison of baseline and follow-up treatment effects suggests that follow-up effects are about 31% as large as the corresponding baseline effects—see [Appendix B.14](#).

In sum, the consumer-harm treatment performs best across intensity, breadth, and durability; the unfair-competition treatment is powerful but narrower; the negative-image treatment has limited breadth and durability; and the market-share treatment has no meaningful effect on demand for antitrust enforcement.

Some additional robustness checks are presented in the Appendix. [Appendix B.10](#) shows that the treatment-effect estimates are robust to adding respondent-level controls and restricting the sample to respondents who passed either the survey attention check or the video-comprehension check. [Appendix B.11](#) reports treatment effects separately for each of the five cases. The main qualitative patterns are broadly evident across the five cases, and none of the results seem to be driven by a single case. [Appendix B.9](#) reports treatment effects separately for the break-up and conduct-remedy components. The component-level estimates closely mirror those for the remedy index. Likewise, [Appendix B.9](#) reports treatment effects separately for each of the six components of the policy index.

The expert forecast survey, analyzed in detail in [Appendix D](#) and summarized below, helps put these results in perspective. Two features of the forecasts are worth emphasizing before comparing them to the experimental estimates. First, experts reported low confidence in their predictions. Second, their forecasts were highly dispersed: for each treatment-outcome pair, some experts predicted negative effects while others predicted sizable positive effects. This lack of consensus suggests that the effects of the information treatments were not obvious *ex ante*. If we compare the experimental estimates to the individual forecasts, only a minority of experts come close to predicting the treatment effects. Even if we average across experts, there are still large discrepancies between the average forecasts and the actual effects. The average forecasts suggest that all treatments would have similar positive effects, whereas the experiment shows large differences across treatments. The contrast is especially stark for market-share information: the treatment experts expected to be the most effective is the least consequential in the experiment.

5 Causal Mechanisms

The previous section shows which information treatments affect demand for antitrust enforcement. This section studies the causal mechanisms behind those effects: which beliefs move respondents toward supporting enforcement and remedies? Because each treatment can shift multiple beliefs at once, we use the experimental variation in belief updating to estimate the relative causal contribution of perceived market share, consumer harm, unfair competition, and negative image.

5.1 2SLS Model

The learning framework in Section 4.1 describes a simple case with one belief and one explicit signal. The mechanism analysis requires a richer framework because our setting has four treatments and four potentially interrelated beliefs. Respondents may use the information from one treatment to make inferences about more than one belief. For example, a respondent who learns that a firm has a higher market share than expected may update the market-share belief directly, but may also infer that the firm causes more consumer harm or view the company more negatively.

Following the econometric approach used in other information-provision experiments, such as Cullen and Perez-Truglia (2022) and Giacobasso et al. (2025), we estimate a 2SLS model that uses the information treatments as exogenous shocks to respondents' belief updates. We retain the notation from Section 4.1, with $K = \{MS, CH, UC, NI\}$ denoting the four belief dimensions and Δb_{ik} denoting respondent i 's belief update along dimension k . The four belief updates are the endogenous variables in the 2SLS model. As in equation (2), the specification uses the prior gap $s_{ic}^{MS} - b_{i,MS}^{prior}$ for market share because that treatment provides a numeric signal. For the other three belief dimensions, it uses prior beliefs, as in equation (3). Let π_c denote case fixed effects for respondent i 's assigned case.

We begin with the first stage. For each belief $k \in K$, we estimate:

$$\begin{aligned} \Delta b_{ik} = & \tau_k + \sum_{\ell \in K} \delta_{k\ell} T_i^\ell + \sum_{\ell \in K} \alpha_{k\ell}^{MS} T_i^\ell (s_{ic}^{MS} - b_{i,MS}^{prior}) \\ & + \sum_{\ell \in K} \sum_{m \in \{CH, UC, NI\}} \alpha_{k\ell m} T_i^\ell b_{im}^{prior} + \eta_k^{MS} (s_{ic}^{MS} - b_{i,MS}^{prior}) \\ & + \sum_{m \in \{CH, UC, NI\}} \eta_{km} b_{im}^{prior} + \pi_c + u_{ik}, \quad k \in K. \end{aligned} \tag{4}$$

Equation (4) is a system of four regressions, one for each of the four beliefs. The treatment indicators allow each information treatment to shift each belief on average. The treatment-

by-prior interactions allow these effects to depend on respondents' starting beliefs. As in Section 4.1, the specification uses the prior gap for the market-share dimension because the market-share signal is explicit; for the other three dimensions, it uses the prior belief itself. It is important to note that this specification does not require each treatment to be a pure shock to only one belief. Instead, it allows, for example, the consumer-harm treatment to affect perceived consumer harm directly while also affecting perceived market share, perceived unfair competition, and perceived negative image.

Let Y_i denote the outcome of interest. The second-stage equation is:

$$Y_i = \omega_0 + \sum_{k \in K} \psi_k \Delta b_{ik} + \gamma_{MS} (s_{ic}^{MS} - b_{i,MS}^{prior}) + \sum_{m \in \{CH, UC, NI\}} \gamma_m b_{im}^{prior} + \pi_c + \epsilon_i, \quad (5)$$

The coefficients ψ_k capture the effect of belief updating along dimension k on the outcome, holding fixed the other belief updates, the market-share prior gap, and the prior levels of the other beliefs. Equation (4) is the corresponding first-stage system for the four endogenous belief updates, Δb_{ik} , in equation (5). The belief updates may be correlated with unobserved determinants of antitrust demand, such as attention to the survey, prior familiarity with the firm, or latent preferences for antitrust enforcement. We therefore estimate equation (5) using a 2SLS model.

The excluded instruments are the treatment indicators, T_i^ℓ , the interactions between the treatment indicators and the market-share prior gap, $T_i^\ell \times (s_{ic}^{MS} - b_{i,MS}^{prior})$, and the interactions between the treatment indicators and the other prior beliefs, $T_i^\ell \times b_{im}^{prior}$, for $\ell \in K$ and $m \in \{CH, UC, NI\}$.³⁰ The treatment indicators isolate average shifts in belief updating induced by the information treatments, while the interactions with prior beliefs and the market-share prior gap allow treatment-induced updating to depend flexibly on respondents' starting beliefs.

5.2 2SLS Results

Before turning to the 2SLS estimates, it is useful to provide an intuition for the identification strategy. The average treatment effects already provide *some* suggestive evidence about the underlying mechanisms. The market-share treatment primarily shifts perceived market share, with much smaller spillovers to the other beliefs. Yet this treatment has little effect on demand for antitrust enforcement. This pattern already suggests that perceived market

³⁰Including the prior beliefs and market-share prior gap as controls ensures that the excluded instruments isolate variation driven by random assignment to the information treatments, rather than cross-sectional differences in respondents' baseline views.

share is unlikely to be an important independent driver of antitrust demand. For the other treatments, however, the average treatment effects are harder to interpret: the consumer-harm, unfair-competition, and negative-image treatments each shift several beliefs at once, so their effects on outcomes could reflect different combinations of perceived consumer harm, perceived unfair competition, and perceived negative image.

The 2SLS strategy builds on the fact that even though the treatments often move several beliefs at once, they do not move them in the same proportions. For instance, the consumer-harm treatment moves all four beliefs, but its largest effect is on perceived consumer harm; the unfair-competition treatment has its largest effect on perceived unfair competition; and the negative-image treatment has its largest effect on perceived negative image. These differences in the relative strength of belief updating help disentangle which belief updates are most closely associated with changes in demand for antitrust enforcement.

Crucially, the identifying variation becomes significantly richer once we account for differences in prior beliefs. Respondents do not update equally in response to the same information: they update more when the information is farther from their priors and less when their priors are already close to the message. As a result, the profile of belief updating induced by a given treatment differs across respondents. For example, a respondent who initially perceives little unfair competition but already perceives substantial consumer harm may update strongly on unfair competition after receiving the unfair-competition treatment, while updating little on consumer harm. The treatment-by-prior interactions in the first stage capture this heterogeneity and help separate the causal contribution of each belief.

To illustrate the richness of the heterogeneous updating by prior beliefs, Figure 4 shows the first-stage results in binned scatterplot form. The diagonal panels (A, F, K, and P), show the effects of each treatment on the main belief it was designed to affect.³¹ The off-diagonal panels, correspond to the spillovers to the other three beliefs.

Figure 4 shows that belief updating is strongly shaped by respondents’ prior beliefs. The reported “Diff. p-value” tests whether the treatment and control slopes are equal. In 13 of the 16 treatment-belief pairs, we reject equality of slopes at conventional levels, indicating that treatment-induced updating usually differs systematically with respondents’ prior beliefs. This pattern appears not only in the diagonal panels, but also in many off-diagonal panels. For example, consider a respondent who initially scores 0 on perceived consumer harm but 4 on perceived unfair competition. In response to the consumer-harm treatment, this respondent updates strongly on perceived consumer harm (Panel F) but does not update at all on perceived unfair competition (Panel K). A respondent with this

³¹Panels A, F, K, and P of Figure 4 are the multivariate counterparts to Panels B, D, F, and H of Figure 3. They need not match exactly because while Figure 3 measures each treatment-belief relationship separately, Figure 4 estimates all treatment-by-prior interactions simultaneously.

combination of prior beliefs is then extremely useful for disentangling the effects of perceived consumer harm from those of perceived unfair competition. The same logic extends to other combinations of prior beliefs: they generate different profiles of belief updating, which help separate the causal contribution of each belief.

Table 4 reports the 2SLS estimates. Panel A reports raw coefficients, but we focus on Panel B, which reports standardized coefficients. The coefficients can be interpreted as the causal effect of a one-point increase in a belief, holding the other three beliefs fixed.

A common concern in 2SLS models is weak-instrument bias (Stock et al., 2002). In our setting, this concern is mitigated by the strong belief updating documented in Section 4.1: the information treatments generate large and statistically significant movements in the belief measures. At the same time, weak instruments are a relevant concern because the four belief updates are interrelated, so the first stage must isolate variation in each belief update holding the others fixed. As a formal diagnostic, the bottom of Table 4 reports the Cragg-Donald Wald F-statistic. The statistic is 9.63 across specifications, close to the conventional rule-of-thumb threshold of 10 (Staiger and Stock, 1997); therefore, we do not view weak instruments as a significant source of concern.

The results from Table 4 indicate that perceived market share has no meaningful effect on demand for antitrust. The effects are close to zero across all four outcomes: -0.003 SD for plaintiff support, -0.004 SD for the remedy index, -0.003 SD for general antitrust support, and 0.002 SD for the policy index. Thus, although the market-share treatment strongly shifts perceived market share, the 2SLS results confirm that perceived market share itself is not an important driver of demand for antitrust enforcement.

Perceived consumer harm is the broadest mechanism. A one-point increase in perceived consumer harm raises plaintiff support by 0.24 SD and the remedy index by 0.18 SD. When respondents learn about the presence of consumer harm, they become more supportive not only of enforcement in the assigned case, but also of broader antitrust policy. It also increases general antitrust support by 0.17 SD and the policy index by 0.33 SD. This is the only belief channel with economically and statistically significant effects across all four outcomes.

Perceived unfair competition is also an important mechanism, but its effects are concentrated on the assigned case. A one-point increase in perceived unfair competition raises plaintiff support by 0.38 SD and the remedy index by 0.35 SD. However, the corresponding effects on general antitrust support and the policy index are small and statistically insignificant. This suggests that perceived unfair conduct makes respondents want enforcement and remedies against the specific firm in the assigned case, but does not generalize as strongly to broader antitrust policy preferences.

As a belief channel, perceived negative image has a narrower effect. A one-point increase

in perceived negative image raises plaintiff support by 0.27 SD, but has little effect on the remedy index, general antitrust support, or the policy index. This evidence suggests that a negative view of the company is enough to move respondents toward the plaintiff side, but not strong enough to prompt support for remedies. This helps interpret the average effect of the negative-image treatment on remedy support discussed above. That treatment shifted not only perceived negative image, but also perceived consumer harm and perceived unfair competition. The 2SLS estimates suggest that the remedy response is more likely driven by the spillover to the beliefs on consumer harm and unfairness than by perceived negative image itself.

The limited role of perceived negative image is also consistent with evidence from the emotions survey. As discussed in Appendix C, all four information treatments elicited substantially more negative than positive emotions, but the emotional response was broadly similar across treatments. This contrasts with the main experiment, where the treatments generated sharply different effects on demand for antitrust enforcement. This pattern suggests that negative affect alone is unlikely to be the main mechanism; instead, the results point to differences in how respondents update specific beliefs and translate those beliefs into policy preferences.

Taken together, the 2SLS estimates suggest a clear hierarchy of mechanisms. Perceived market share is not important. Perceived consumer harm is the broadest and most robust channel. Perceived unfair competition is powerful for plaintiff support and remedies, while perceived negative image mainly affects plaintiff support. The expert forecasts also missed the mark on causal mechanism. In particular, experts predicted that perceived market share would be by far the most important belief channel, whereas the 2SLS results indicate that it does not matter at all—see Appendix D for more details.

6 Conclusions

This paper studies public demand for antitrust enforcement using a pre-registered survey experiment with 4,016 American households. Respondents were assigned to one of five real antitrust cases and then randomly assigned to receive no additional information or one of four information treatments: market share, consumer harm, unfair competition, or negative image. The treatments had sharply different effects on demand for antitrust enforcement. Information about market share substantially changed perceived market share but did not meaningfully increase plaintiff support or remedy support. The other treatments increased antitrust demand, but differed in their intensity, breadth, and durability. The consumer-harm treatment had the strongest profile on all three dimensions: it increased plaintiff support and

support for consequential remedies, spilled over to broader antitrust policy preferences, and persisted in the follow-up survey. The unfair-competition treatment also increased plaintiff support and remedy support, but its effects were less broad and less durable. The negative-image treatment had the narrowest effects: it shifted respondents toward the plaintiff side and had smaller effects on remedies, but generated little breadth or durability.

Because each treatment can move several beliefs at once, these average treatment effects do not by themselves reveal which beliefs causally drive demand for antitrust enforcement. We therefore estimate a 2SLS model to disentangle the causal effect of each of the four beliefs. The results point to a clear hierarchy of mechanisms. Perceived consumer harm is the most robust driver of antitrust demand, increasing support for the plaintiff side, remedies, and even preferences about dominant firms beyond the specific case at hand. Perceived unfair competition matters for plaintiff support and remedies but lacks breadth. Perceived negative image plays a more limited role, while perceived market share has no meaningful effects.

Our results suggest that public demand for antitrust aligns closely with the standard economic framework in its sharp distinction between *market share* and *market power*: respondents do not demand antitrust enforcement simply because a firm has a high market share, but only if they believe there is consumer harm. The departure from the core economic framework comes from the perceived unfair competition. Holding perceived consumer harm fixed, respondents are more willing to support enforcement and remedies when they believe dominance was obtained or maintained through unfair practices, although this channel is more limited in breadth and durability.

The fact that the public cares about consumer harm and not market share per se is perhaps most surprising given the prominent role that market share often plays in antitrust litigation, typically in the form of market definition. For example, in the FTC’s case against Meta, the FTC alleged that Meta had monopoly power in a market for “personal social networking” services, but the court rejected this narrow market definition, emphasizing that the relevant competitive landscape also had to account for substitutes such as TikTok and YouTube. Once the market was defined more broadly, Meta’s share was below the level typically associated with monopoly power, and the court ruled for Meta ([U.S. District Court for the District of Columbia, 2025](#)). Moreover, the discussion of market share was also prominent in the other four cases included in our study.³² This illustrates how, in legal practice, the assessment of market power can hinge on the antecedent question of market definition: the same firm may

³²For Meta, see [U.S. District Court for the District of Columbia \(2025, pp. 82–87\)](#). For Google, see [U.S. District Court for the District of Columbia \(2024, pp. 156–157\)](#); for Live Nation, see [U.S. District Court for the Southern District of New York \(2026, p. 36\)](#); for Apple/Epic, see [U.S. Court of Appeals for the Ninth Circuit \(2023, p. 41\)](#); for Luxottica, see [U.S. District Court for the Southern District of New York \(2025, pp. 20–26\)](#).

appear dominant in a narrowly defined market but substantially less so in a broader market. In this sense, public preferences line up closely with the economic framework instead of the market-share-centered approach often emphasized in legal practice.

Our findings also have implications for how policymakers and regulators communicate about antitrust enforcement. Public discussions of antitrust often emphasize the sheer size of dominant firms, including their market shares and the concentration of the markets in which they operate. Our results suggest that this type of information is unlikely, on its own, to increase public support for enforcement. Communication strategies that focus instead on prices, quality, choice, and other consumer-welfare consequences may be more effective at generating durable support for remedial measures, such as breaking up companies, and may also increase broader support for antitrust policy. Persuading the public that large firms engage in unfair practices may also increase support, especially when the goal is to build demand for enforcement or remedies in a specific case. However, these fairness-based arguments seem less likely to generate broader support for antitrust policy.

References

- Affonso, Felipe (2026) “Brief Commentary: A Framework for Detecting AI Agents in Online Research,” *Journal of Consumer Research*, ucag006.
- Autor, David, David Dorn, Lawrence F. Katz, Christina Patterson, and John Van Reenen (2020) “The Fall of the Labor Share and the Rise of Superstar Firms,” *The Quarterly Journal of Economics*, 135 (2), 645–709.
- Axios (2026) “Meta Outspends Big Tech Peers on Lobbying Again,” *January 21, 2026*.
- Baker, Richard B., Carola Frydman, and Eric Hilt (2023) “Political Discretion and Antitrust Policy: Evidence from the Assassination of President McKinley,” *Journal of Law and Economics*, 66 (4), 837–873.
- Brutger, Ryan and Amy Pond (2025a) “International Economic Relations and American Support for Antitrust Policy,” *Business and Politics*, 27 (2), 159–179.
- (2025b) “Partisan Preferences for Antitrust Policy,” *Political Science Research and Methods*, 1–11.
- Cavallo, Alberto, Guillermo Cruces, and Ricardo Perez-Truglia (2017) “Inflation expectations, learning, and supermarket prices: Evidence from survey experiments,” *American Economic Journal: Macroeconomics*, 9 (3), 1–35.
- Çelebi, Can, Christine Exley, Sören Harrs, Hannu Kivimaki, Marta Serra-Garcia, and Jeffrey Yusof (2026) “Mission Possible: The Collection of High-Quality Online Data,” *CESifo Working Paper No. 12535*.
- Cullen, Zoë B. and Ricardo Perez-Truglia (2022) “How Much Does Your Boss Make? The Effects of Salary Comparisons,” *Journal of Political Economy*, 130 (3), 766–822.
- DeGroot, Morris H. (1970) *Optimal Statistical Decisions*: McGraw-Hill.
- DellaVigna, Stefano, Nicholas Otis, and Eva Vivalt (2020) “Forecasting the Results of Experiments: Piloting an Elicitation Strategy,” *AEA Papers and Proceedings*, 110, 75–79.
- Demsetz, Harold (1973) “Industry Structure, Market Rivalry, and Public Policy,” *Journal of Law and Economics*, 16 (1), 1–9.
- Di Tella, Rafael, Juan Dubra, and Alejandro Lagomarsino (2021) “Meet the Oligarchs: Business Legitimacy and Taxation at the Top,” *The Journal of Law and Economics*, 64 (4), 651–674.
- Dohmen, Thomas and Tomáš Jagelka (2024) “Accounting for individual-specific reliability of self-assessed measures of economic preferences and personality traits,” *Journal of Political Economy Microeconomics*, 2 (3), 399–462.

- Dreber, Anna, Thomas Pfeiffer, Johan Almenberg, Siri Isaksson, Brad Wilson, Yiling Chen, Brian A Nosek, and Magnus Johannesson (2015) “Using prediction markets to estimate the reproducibility of scientific research,” *Proceedings of the National Academy of Sciences*, 112 (50), 15343–15347.
- Elías, Julio J, Nicola Lacetera, and Mario Macis (2019) “Paying for kidneys? A randomized survey and choice experiment,” *American Economic Review*, 109 (8), 2855–2888.
- Farrell, Joseph and Paul Klemperer (2007) “Coordination and Lock-In: Competition with Switching Costs and Network Effects,” in Armstrong, Mark and Robert Porter eds. *Handbook of Industrial Organization*, 3, 1967–2072: Elsevier.
- Fisman, Raymond, Keith Gladstone, Ilyana Kuziemko, and Suresh Naidu (2020) “Do Americans want to tax wealth? Evidence from online surveys,” *Journal of Public Economics*, 188, 104207.
- Fortune (2026) “Big Tech Is Spending \$226,000 a Day on Lobbying Congress, Advocacy Group Finds,” *April 23, 2026*.
- Giacobasso, Matias, Brad Nathan, Ricardo Perez-Truglia, and Alejandro Zentner (2025) “Where Do My Tax Dollars Go? Tax Morale Effects of Perceived Government Spending,” *American Economic Journal: Applied Economics*, 17 (4), 223–259.
- Hope, David, Julian Limberg, and Nina Weber (2023) “Why do (some) ordinary Americans support tax cuts for the rich? Evidence from a randomised survey experiment,” *European Journal of Political Economy*, 78, 102349.
- Jiang, Chenxi, Maximiliano Lauletta, Ro’ee Levy, Joseph S. Shapiro, and Dmitry Taubinsky (2026) “Understanding Support for Inefficient Environmental Policy Instruments,” *NBER Working Paper No. 35073*.
- Kahneman, Daniel, Jack L. Knetsch, and Richard H. Thaler (1986) “Fairness as a Constraint on Profit Seeking: Entitlements in the Market,” *The American Economic Review*, 76 (4), 728–741.
- Klemperer, Paul (1987) “Markets with Consumer Switching Costs,” *Quarterly Journal of Economics*, 102 (2), 375–394.
- Kling, Jeffrey R., Jeffrey B. Liebman, and Lawrence F. Katz (2007) “Experimental Analysis of Neighborhood Effects,” *Econometrica*, 75 (1), 83–119.
- Krueger, Alan B. and David A. Schkade (2008) “The reliability of subjective well-being measures,” *Journal of Public Economics*, 92 (8), 1833–1845.
- Kuziemko, Ilyana, Michael I. Norton, Emmanuel Saez, and Stefanie Stantcheva (2015) “How

- Elastic Are Preferences for Redistribution? Evidence from Randomized Survey Experiments,” *American Economic Review*, 105 (4), 1478–1508.
- Kwoka, John E. and Diana L. Moss (2012) “Behavioral Merger Remedies: Evaluation and Implications for Antitrust Enforcement,” *The Antitrust Bulletin*, 57 (4), 979–1011.
- Lancieri, Filippo, Eric A. Posner, and Luigi Zingales (2023) “The Political Economy of the Decline of Antitrust Enforcement in the United States,” *Antitrust Law Journal*, 85 (2), 441–519.
- New York Times (2024) “How the Google Antitrust Ruling May Influence Tech Competition,” *August 12, 2024*.
- Perez-Truglia, Ricardo and Jeffrey Yusof (2026) “Billionaire Superstar: Public Image and Demand for Taxation,” NBER Working Paper No. 32712.
- Prolific (2026) “How Prolific Detects Bots and AI in Online Research,” <https://www.prolific.com/resources/how-prolific-detects-bots-and-ai-in-online-research>.
- Sapienza, Paola and Luigi Zingales (2013) “Economic Experts versus Average Americans,” *American Economic Review*, 103 (3), 636–642.
- Shapiro, Carl (2019) “Protecting Competition in the American Economy: Merger Control, Tech Titans, Labor Markets,” *Journal of Economic Perspectives*, 33 (3), 69–93.
- Staiger, Douglas and James H. Stock (1997) “Instrumental Variables Regression with Weak Instruments,” *Econometrica*, 65 (3), 557–586.
- Stantcheva, Stefanie (2021) “Understanding tax policy: How do people reason?,” *The Quarterly Journal of Economics*, 136 (4), 2309–2369.
- Stock, James H., Jonathan H. Wright, and Motohiro Yogo (2002) “A Survey of Weak Instruments and Weak Identification in Generalized Method of Moments,” *Journal of Business & Economic Statistics*, 20 (4), 518–529.
- Tirole, Jean (1988) *The Theory of Industrial Organization*, Cambridge, MA: MIT Press.
- U.S. Court of Appeals for the Ninth Circuit (2023) “Epic Games, Inc. v. Apple, Inc.,” No. 21-16506; 67 F.4th 946.
- (2025) “Epic Games, Inc. v. Google LLC et al. (In re Google Play Store Antitrust Litigation),” Nos. 24-6256, 24-6274, 25-303, 147 F.4th 917.
- U.S. District Court for the District of Columbia (2024) “United States v. Google LLC, Memorandum Opinion,” Case Nos. 20-cv-3010 (APM) and 20-cv-3715 (APM); 2024 WL 3647498.
- (2025) “Federal Trade Commission v. Meta Platforms, Inc., Redacted Memorandum

Opinion,” No. 1:20-cv-03590 (JEB), Document 693.

U.S. District Court for the Southern District of New York (2025) “In re Eyewear Antitrust Litigation, Opinion and Order Granting Motion to Dismiss,” No. 1:24-cv-4826 (MKV), Document 250.

——— (2026) “United States v. Live Nation Entertainment, Inc. and Ticketmaster L.L.C., Opinion and Order,” No. 24-cv-3973 (AS), Document 1037.

U.S. Supreme Court (1959) “Beacon Theatres, Inc. v. Westover,” 359 U.S. 500.

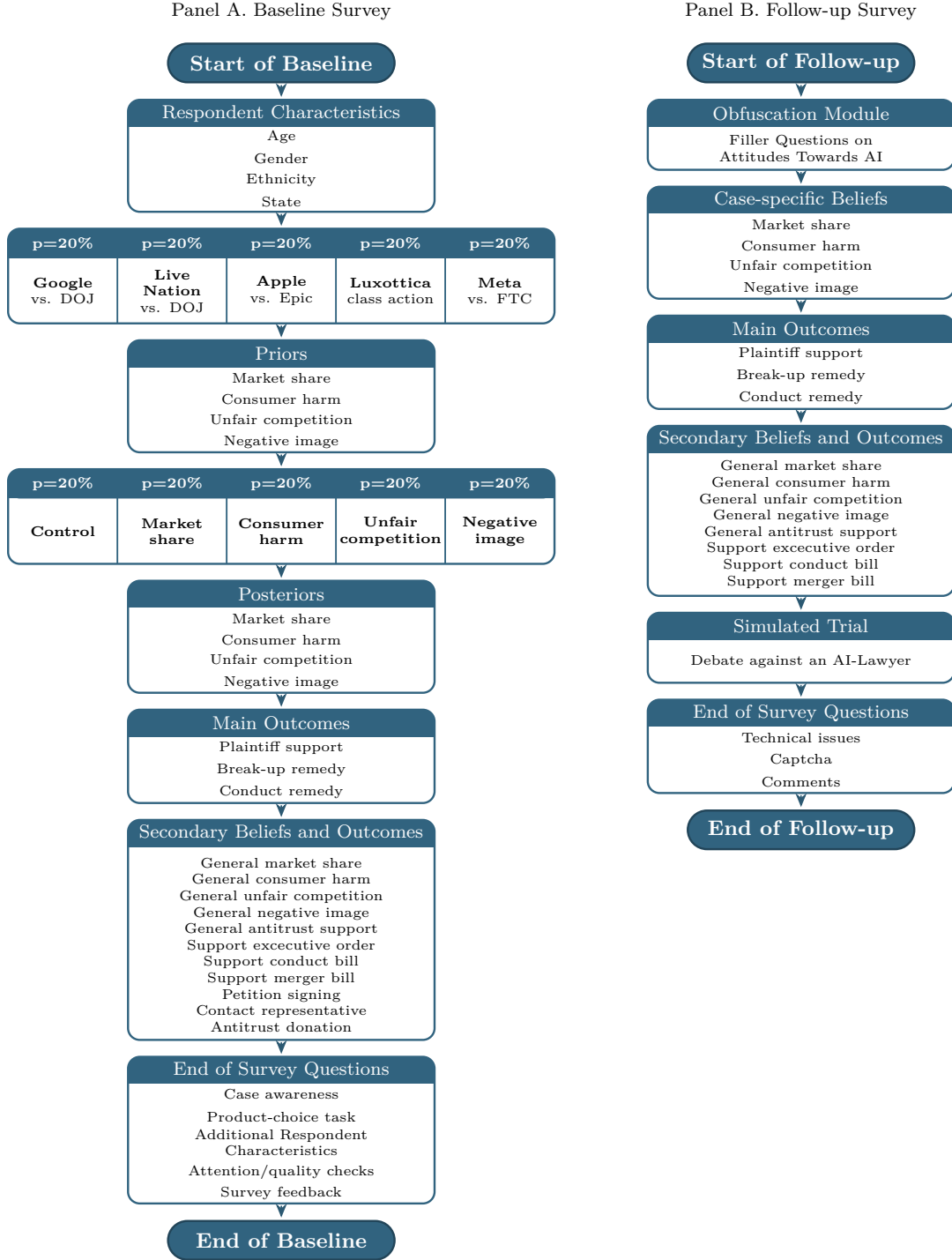
Wall Street Journal (2026) “The AI-Fication of the S&P 500,” *The Wall Street Journal*, June 1 2026.

Washington Post (2024a) “Elon Musk Is Now America’s Largest Political Donor,” *December 6, 2024*.

——— (2024b) “Where All the Biden Administration’s Big Tech Antitrust Lawsuits Stand,” *October 9, 2024*.

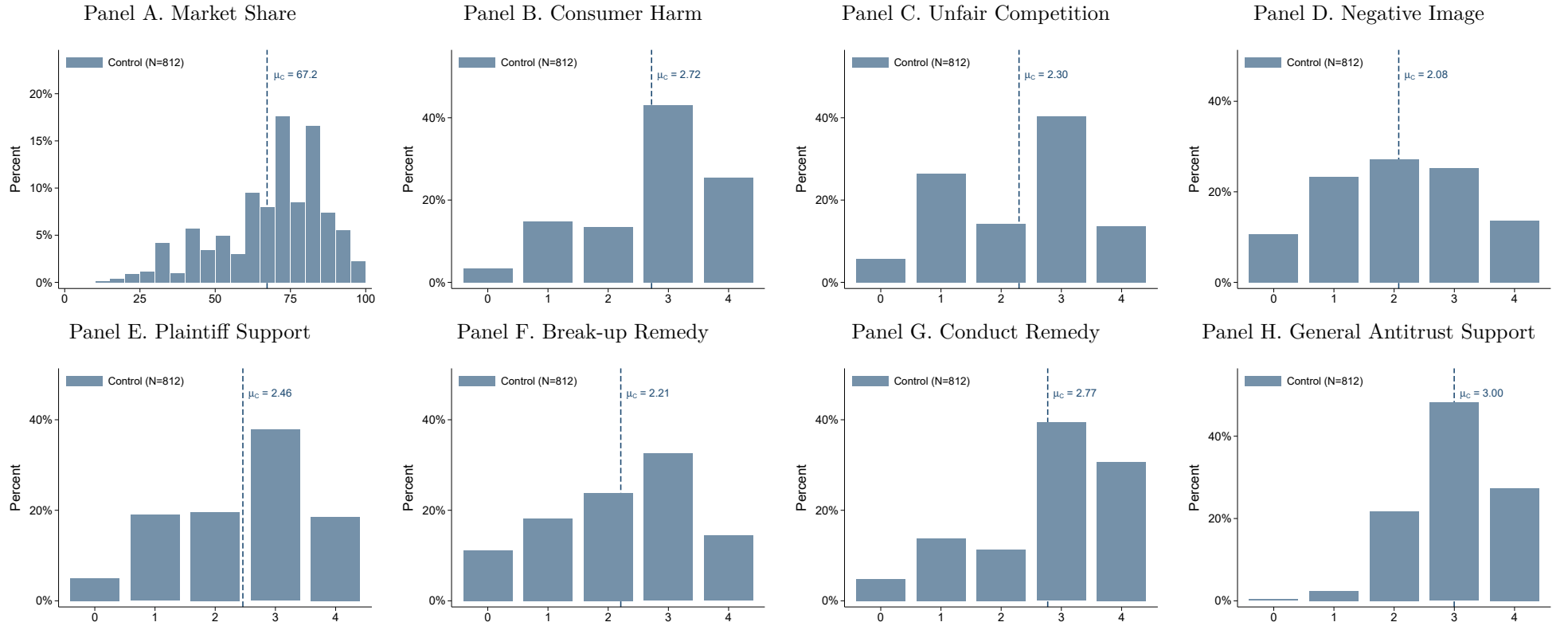
Wolfe, Jan (2024) “Big Tech Braces for Wave of Antitrust Rulings in 2024,” *The Wall Street Journal*, January 1, 2024.

Figure 1: Survey Flow



Notes: The figure summarizes the sequence of the baseline survey (Panel A) and follow-up survey (Panel B). At baseline, respondents were assigned with equal probability to one of five antitrust cases and subsequently to the control group or one of four case-specific information treatments. Posterior beliefs repeat the four prior-belief measures. The follow-up was administered approximately one month later and repeated the belief measures and most outcomes, but not petition signing, contacting a representative, or the donation-allocation task. The trial simulation followed all repeated measures.

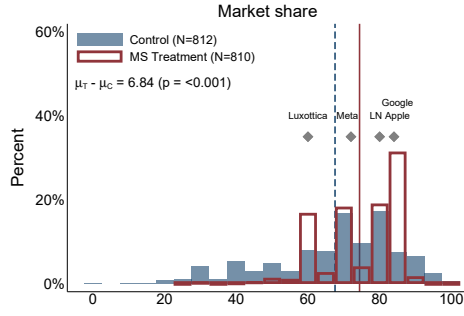
Figure 2: Baseline Preferences and Beliefs in the Control Group



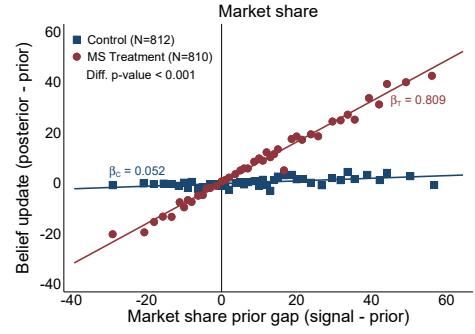
Notes: The figure reports distributions for respondents assigned to the control group. Panels A–D report prior beliefs about the assigned case. Panels E–G show case-specific outcomes, and Panel H shows general support for antitrust enforcement. The market-share belief is measured on a 0 to 100 percent. All other outcomes range from 0 to 4. Higher values indicate greater perceived consumer harm, more unfair competition, a more negative company image, stronger plaintiff or remedy support, and stricter preferred antitrust enforcement, respectively. Dashed lines and μ_C denote panel-specific control-group means; legends report the number of observations.

Figure 3: Treatment Effects on Beliefs

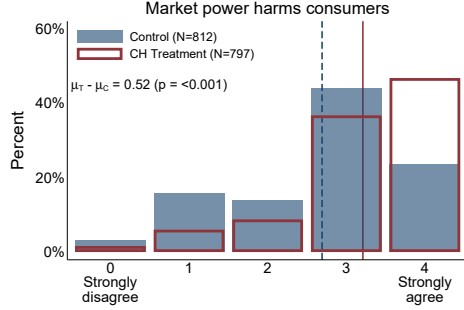
Panel A. Market Share: Posterior Beliefs



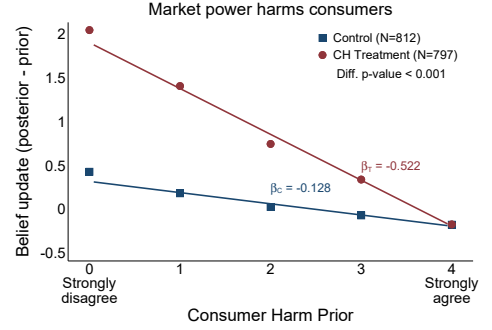
Panel B. Market Share: Belief Updating



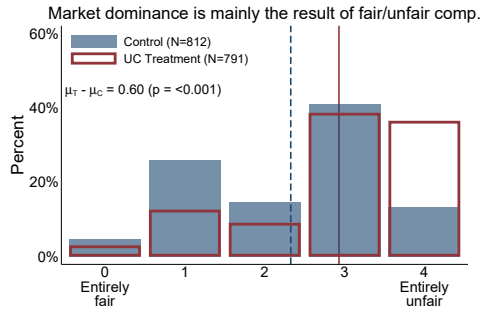
Panel C. Consumer Harm: Posterior Beliefs



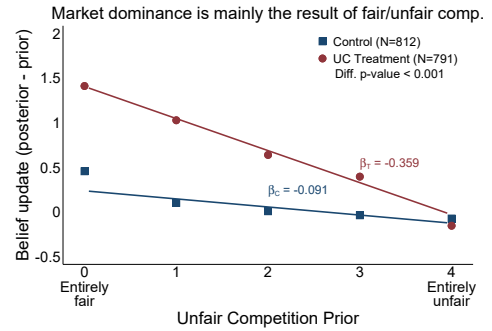
Panel D. Consumer Harm: Belief Updating



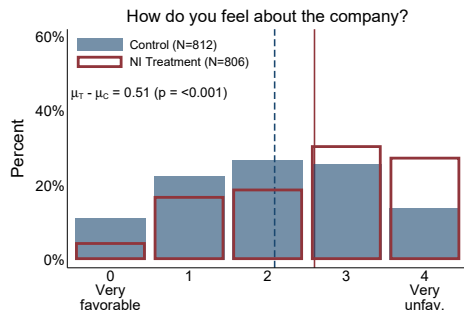
Panel E. Unfair Competition: Posterior Beliefs



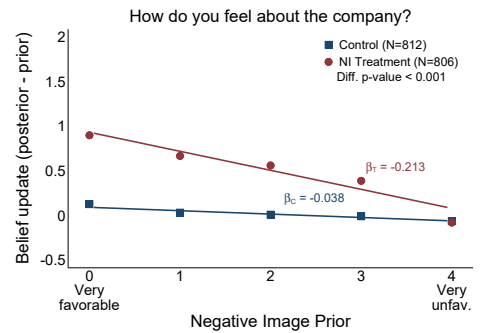
Panel F. Unfair Competition: Belief Updating



Panel G. Negative Image: Posterior Beliefs

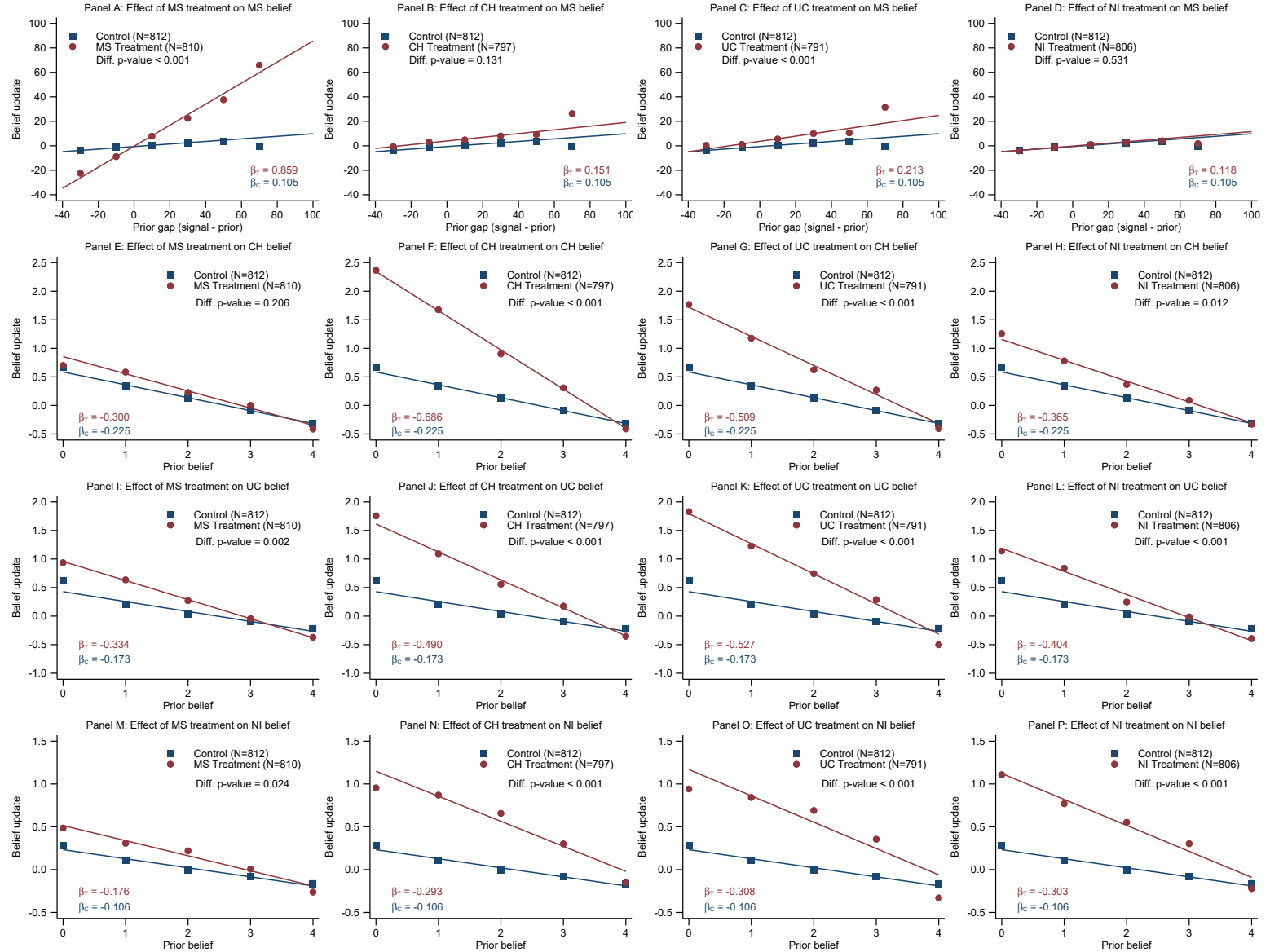


Panel H. Negative Image: Belief Updating



Notes: The figure compares beliefs between respondents assigned to each information treatment and respondents assigned to the control group. The left panels show posterior-belief distributions. The right panels show belief updating, defined as the posterior minus the prior belief, as a function of prior beliefs. For market share, the horizontal axis is the prior gap, defined as the case-specific signal minus the prior belief. The other beliefs range from 0 to 4. In the left panels, $\mu_T - \mu_C$ denotes the difference in treatment- and control-group means; the corresponding p -value tests equality of means. In the right panels, β_T and β_C denote the treatment- and control-group slopes; the reported p -value tests equality of slopes. Legends report the number of observations.

Figure 4: First-Stage Belief Updating



Notes: The figure visualizes the first-stage belief-updating regressions in equation (4). Rows correspond to the belief being updated, and columns correspond to the information treatment. MS, CH, UC, and NI denote market share, consumer harm, unfair competition, and negative image, respectively. Belief updating is defined as the posterior minus the prior belief. Lines show fitted values from regressions of belief updating on treatment assignment, prior beliefs, treatment-by-prior interactions, and assigned-case fixed effects. Markers show regression-adjusted means for each prior-belief category: 20-point intervals for the market-share prior gap and response categories 0–4 for the other priors. The other prior beliefs are held at their sample means, and all regressions include assigned-case fixed effects. In the market-share row, the horizontal axis is the prior gap, defined as the case-specific signal minus the prior belief; positive values indicate prior underestimation. The other rows use prior beliefs measured from 0 to 4. The labels β_C and β_T denote the control- and treatment-group slopes. “Diff. p-value” tests whether these slopes are equal. Legends report the number of observations.

Table 1: Descriptive Statistics and Balance Tests

	Full Sample (1)	By Treatment Group					p-value (7)
		Control (2)	Market Share (3)	Consumer Harm (4)	Unfair Competition (5)	Negative Image (6)	
<i>Panel A: Respondent characteristics</i>							
Female (%)	57.7 (0.8)	57.9 (1.7)	58.1 (1.7)	57.0 (1.8)	58.4 (1.8)	57.3 (1.7)	0.977
Age ≤ 35 (%)	45.9 (0.8)	46.6 (1.8)	44.8 (1.7)	44.0 (1.8)	46.6 (1.8)	47.6 (1.8)	0.593
White (%)	66.0 (0.7)	66.9 (1.7)	66.2 (1.7)	64.1 (1.7)	66.1 (1.7)	66.6 (1.7)	0.792
Black (%)	18.6 (0.6)	19.3 (1.4)	18.8 (1.4)	19.4 (1.4)	16.2 (1.3)	19.4 (1.4)	0.402
Asian (%)	5.6 (0.4)	4.9 (0.8)	5.9 (0.8)	6.1 (0.9)	6.3 (0.9)	4.6 (0.7)	0.455
Hispanic (%)	6.3 (0.4)	5.5 (0.8)	6.3 (0.9)	7.2 (0.9)	7.5 (0.9)	5.2 (0.8)	0.271
Income ≥ 50k (%)	65.6 (0.7)	64.5 (1.7)	66.8 (1.7)	65.5 (1.7)	66.6 (1.7)	64.5 (1.7)	0.789
College degree (%)	54.8 (0.8)	53.3 (1.8)	55.8 (1.7)	56.0 (1.8)	52.5 (1.8)	56.6 (1.7)	0.375
Democrat (%)	43.4 (0.8)	44.2 (1.7)	45.7 (1.8)	42.8 (1.8)	40.8 (1.7)	43.4 (1.7)	0.383
Republican (%)	25.0 (0.7)	25.1 (1.5)	24.7 (1.5)	23.1 (1.5)	27.3 (1.6)	24.9 (1.5)	0.427
Trust in federal government	1.00 (0.01)	1.00 (0.03)	1.02 (0.03)	0.99 (0.03)	1.00 (0.03)	1.00 (0.02)	0.958
<i>Panel B: Follow-up participation</i>							
Completed follow-up survey (%)	71.5 (0.7)	73.2 (1.6)	70.4 (1.6)	72.5 (1.6)	70.9 (1.6)	70.3 (1.6)	0.618
N	4,016	812	810	797	791	806	

Notes: Panel A reports baseline respondent characteristics, and Panel B reports the share of baseline respondents who completed the follow-up survey. Column (1) reports means for the full baseline sample. Columns (2)–(6) report means by randomized information-treatment assignment. Column (7) reports the p -value from a joint test that the means are equal across the five treatment groups, conducted separately for each row. Standard errors are reported in parentheses. Indicator variables are reported in percentage points. Trust in the federal government ranges from 0 to 3, with higher values indicating greater trust.

Table 2: Average Treatment Effects on Case-Specific Beliefs and Main Outcomes

	Baseline Survey						Follow-Up Survey					
	Beliefs				Outcomes		Beliefs				Outcomes	
	(1) Mkt. share	(2) Cons. harm	(3) Unfair	(4) Neg. image	(5) Plaintiff supp.	(6) Remedy index	(7) Mkt. share	(8) Cons. harm	(9) Unfair	(10) Neg. image	(11) Plaintiff supp.	(12) Remedy index
<i>Panel A: Raw Coefficients</i>												
Market share treatment	6.840*** (0.732)	0.040 (0.055)	0.143** (0.057)	0.078 (0.062)	0.102* (0.057)	0.057 (0.051)	1.730* (0.956)	-0.021 (0.061)	-0.025 (0.065)	-0.096 (0.071)	0.008 (0.067)	0.021 (0.060)
Consumer harm treatment	5.960*** (0.866)	0.519*** (0.051)	0.467*** (0.057)	0.539*** (0.061)	0.512*** (0.054)	0.291*** (0.050)	2.500** (0.980)	0.111* (0.059)	0.064 (0.065)	0.153** (0.069)	0.168*** (0.064)	0.131** (0.059)
Unfair competition treatment	6.377*** (0.859)	0.408*** (0.053)	0.604*** (0.056)	0.533*** (0.062)	0.544*** (0.054)	0.342*** (0.049)	1.703* (1.021)	0.068 (0.061)	0.155** (0.065)	0.104 (0.071)	0.129* (0.066)	0.104* (0.059)
Negative image treatment	0.343 (0.885)	0.210*** (0.053)	0.241*** (0.056)	0.507*** (0.060)	0.328*** (0.055)	0.200*** (0.051)	0.493 (1.016)	-0.019 (0.060)	0.045 (0.065)	0.066 (0.070)	0.072 (0.065)	0.017 (0.062)
<i>Panel B: Standardized Coefficients</i>												
Market share treatment	0.388*** (0.042)	0.037 (0.050)	0.126** (0.051)	0.064 (0.051)	0.090* (0.050)	0.055 (0.050)	0.098* (0.054)	-0.020 (0.058)	-0.022 (0.058)	-0.080 (0.059)	0.007 (0.058)	0.020 (0.058)
Consumer harm treatment	0.338*** (0.049)	0.477*** (0.047)	0.411*** (0.050)	0.443*** (0.050)	0.448*** (0.047)	0.281*** (0.048)	0.141** (0.055)	0.105* (0.057)	0.057 (0.058)	0.127** (0.057)	0.147*** (0.056)	0.124** (0.056)
Unfair competition treatment	0.362*** (0.049)	0.375*** (0.048)	0.532*** (0.049)	0.439*** (0.051)	0.477*** (0.047)	0.331*** (0.047)	0.096* (0.058)	0.064 (0.058)	0.139** (0.058)	0.087 (0.059)	0.112* (0.058)	0.099* (0.057)
Negative image treatment	0.019 (0.050)	0.193*** (0.048)	0.212*** (0.049)	0.417*** (0.049)	0.288*** (0.048)	0.193*** (0.049)	0.028 (0.057)	-0.018 (0.057)	0.040 (0.058)	0.055 (0.058)	0.063 (0.057)	0.016 (0.059)
Observations	4,016	4,016	4,016	4,016	4,016	4,016	2,869	2,869	2,869	2,869	2,869	2,869
Control Mean	67.547	2.692	2.326	2.089	2.459	2.491	68.891	2.875	2.488	2.259	2.643	2.621
Control SD	17.632	1.087	1.136	1.215	1.142	1.035	17.709	1.051	1.120	1.204	1.146	1.048

Notes: The table reports average treatment effects of the four randomized information treatments relative to the control group. Each column is estimated using a separate OLS regression of the indicated outcome on the four treatment indicators. Columns (1)–(6) report baseline-survey effects; columns (7)–(12) report effects measured in the follow-up survey approximately one month later. Columns (1)–(4) and (7)–(10) report effects on posterior beliefs, while columns (5)–(6) and (11)–(12) report effects on case-specific outcomes. Panel A reports coefficients in the original units. Market share ranges from 0 to 100 percent; all other beliefs and outcomes range from 0 to 4, with higher values indicating greater perceived consumer harm, more unfair competition, a more negative company image, stronger plaintiff support, or stronger remedy support. The remedy index is the mean of the break-up and conduct-remedy outcomes. Panel B standardizes each outcome using its control-group mean and standard deviation separately within each survey. The reported control means and standard deviations are in the original units. Robust standard errors are reported in parentheses. ***, **, and * denote significance at the 1, 5, and 10 percent levels, respectively.

Table 3: Average Treatment Effects on General Beliefs and Additional Outcomes

	Baseline Survey						Follow-Up Survey					
	General Beliefs				Outcomes		General Beliefs				Outcomes	
	(1) Mkt. share	(2) Cons. harm	(3) Unfair	(4) Neg. image	(5) Gen. antitrust	(6) Policy index	(7) Mkt. share	(8) Cons. harm	(9) Unfair	(10) Neg. image	(11) Gen. antitrust	(12) Policy index
<i>Panel A: Raw Coefficients</i>												
Market share treatment	3.015*** (0.900)	0.025 (0.047)	0.028 (0.056)	-0.051 (0.053)	-0.002 (0.040)	-0.035 (0.049)	1.015 (1.017)	-0.060 (0.054)	-0.121* (0.066)	-0.135** (0.061)	-0.004 (0.046)	-0.175*** (0.062)
Consumer harm treatment	2.461*** (0.937)	0.233*** (0.045)	0.195*** (0.055)	0.157*** (0.051)	0.100*** (0.039)	0.053 (0.048)	1.273 (1.015)	-0.012 (0.053)	-0.000 (0.063)	0.041 (0.060)	0.070 (0.045)	-0.128** (0.059)
Unfair competition treatment	1.617* (0.934)	0.154*** (0.046)	0.134** (0.056)	0.084 (0.052)	0.064 (0.040)	-0.050 (0.050)	1.074 (1.043)	-0.015 (0.052)	-0.055 (0.065)	-0.010 (0.060)	0.019 (0.046)	-0.158*** (0.061)
Negative image treatment	0.169 (0.946)	0.085* (0.047)	0.099* (0.056)	0.066 (0.052)	0.064 (0.040)	-0.012 (0.049)	0.071 (1.044)	-0.044 (0.053)	-0.080 (0.065)	-0.046 (0.061)	0.005 (0.047)	-0.106* (0.060)
<i>Panel B: Standardized Coefficients</i>												
Market share treatment	0.156*** (0.047)	0.026 (0.050)	0.025 (0.050)	-0.049 (0.051)	-0.003 (0.050)	-0.035 (0.049)	0.056 (0.056)	-0.068 (0.060)	-0.110* (0.059)	-0.132** (0.060)	-0.006 (0.059)	-0.175*** (0.062)
Consumer harm treatment	0.128*** (0.049)	0.247*** (0.048)	0.174*** (0.049)	0.151*** (0.049)	0.128*** (0.049)	0.053 (0.048)	0.070 (0.056)	-0.013 (0.060)	-0.000 (0.057)	0.041 (0.059)	0.089 (0.058)	-0.128** (0.059)
Unfair competition treatment	0.084* (0.048)	0.163*** (0.048)	0.120** (0.050)	0.081 (0.050)	0.082 (0.050)	-0.050 (0.050)	0.059 (0.058)	-0.016 (0.058)	-0.049 (0.059)	-0.010 (0.059)	0.024 (0.059)	-0.158*** (0.061)
Negative image treatment	0.009 (0.049)	0.090* (0.050)	0.089* (0.050)	0.064 (0.050)	0.082 (0.051)	-0.012 (0.049)	0.004 (0.058)	-0.049 (0.059)	-0.072 (0.059)	-0.045 (0.060)	0.006 (0.060)	-0.106* (0.060)
Observations	4,016	4,016	4,016	4,016	4,016	4,016	2,869	2,869	2,869	2,869	2,869	2,869
Control Mean	64.161	3.021	2.344	2.416	2.998	0.000	65.980	3.146	2.493	2.495	3.022	0.000
Control SD	19.292	0.941	1.121	1.040	0.787	1.000	18.079	0.896	1.107	1.021	0.782	1.000

Notes: The table reports average treatment effects of the four randomized information treatments relative to the control group. Each column is estimated using a separate OLS regression of the indicated outcome on the four treatment indicators. Columns (1)–(6) report baseline-survey effects; columns (7)–(12) report effects measured in the follow-up survey approximately one month later. Columns (1)–(4) and (7)–(10) report effects on posterior beliefs about dominant firms generally. Columns (5) and (11) report effects on general support for antitrust enforcement, while columns (6) and (12) report effects on the policy index. Panel A reports coefficients in the original units. General market share ranges from 0 to 100 percent; the other general beliefs and general antitrust support range from 0 to 4, with higher values indicating greater perceived consumer harm, more unfair competition, a more negative image of dominant firms, or stricter preferred antitrust enforcement. The policy index is constructed following [Kling et al. \(2007\)](#), with higher values indicating greater policy support. At baseline, it combines three policy-support questions, petition signing, contacting a representative, and the antitrust-donation allocation. At follow-up, it combines only the three policy-support questions because the other components were not re-elicited. Panel B standardizes every outcome using its control-group mean and standard deviation separately within each survey; therefore, the policy-index coefficients are unchanged across panels. The reported control means and standard deviations are in the original units. Robust standard errors are reported in parentheses. ***, **, and * denote significance at the 1, 5, and 10 percent levels, respectively.

Table 4: 2SLS Effects of Belief Updating on Antitrust Outcomes

	(1) Plaintiff support	(2) Remedy index	(3) General antitrust	(4) Policy index
<i>Panel A: Raw Coefficients</i>				
Perceived Market Share	-0.003 (0.002)	-0.004* (0.002)	-0.003 (0.002)	0.002 (0.003)
Perceived Consumer Harm	0.269*** (0.091)	0.174** (0.082)	0.136* (0.076)	0.328*** (0.097)
Perceived Unfair Competition	0.433*** (0.105)	0.345*** (0.098)	0.052 (0.082)	-0.157 (0.106)
Perceived Negative Image	0.303*** (0.094)	0.079 (0.089)	0.007 (0.078)	-0.147 (0.102)
<i>Panel B: Standardized Coefficients</i>				
Perceived Market Share	-0.003 (0.002)	-0.004* (0.002)	-0.003 (0.002)	0.002 (0.003)
Perceived Consumer Harm	0.235*** (0.080)	0.168** (0.079)	0.173* (0.097)	0.328*** (0.097)
Perceived Unfair Competition	0.379*** (0.092)	0.333*** (0.095)	0.066 (0.104)	-0.157 (0.106)
Perceived Negative Image	0.265*** (0.083)	0.076 (0.086)	0.009 (0.099)	-0.147 (0.102)
Observations	4,016	4,016	4,016	4,016
Baseline Mean	2.459	2.491	2.998	0.000
Baseline SD	1.142	1.035	0.787	1.000
Cragg-Donald Wald F-statistic	9.63	9.63	9.63	9.63

Notes: The table reports 2SLS estimates using baseline-survey responses. Each column is estimated using a separate regression. The four endogenous regressors are the belief updates in perceived market share, consumer harm, unfair competition, and negative image, each defined as the posterior minus the prior belief. The excluded instruments are the four treatment indicators and their interactions with each of the four prior beliefs. For market share, the belief measure is the prior gap, defined as the case-specific signal minus the prior belief; the other prior beliefs are measured on their original 0–4 scales. All specifications include the four belief measures and assigned-case fixed effects. Columns (1) and (2) report effects on case-specific plaintiff and remedy support; columns (3) and (4) report effects on general antitrust support and the policy index. Panel A reports coefficients in the original outcome units. Plaintiff support, the remedy index, and general antitrust support range from 0 to 4. The remedy index is the mean of the break-up and conduct-remedy outcomes. The policy index combines the three policy-support questions, petition signing, contacting a representative, and the antitrust-donation allocation following Kling et al. (2007); it is standardized using the control group. Panel B standardizes the dependent variable using its control-group mean and standard deviation but leaves the belief updates in their original units. The reported baseline means and standard deviations are control-group values in the original units. The table also reports the first-stage Cragg–Donald Wald F -statistic. Robust standard errors are reported in parentheses. ***, **, and * denote significance at the 1, 5, and 10 percent levels, respectively.