

Sources of Market Power in Web Search: Evidence from a Field Experiment

Hunt Allcott, Juan Camilo Castillo, Matthew Gentzkow, Leon Musolff, and Tobias Salz*

May 13, 2025

Abstract

We study the forces behind Google’s large web-search market share. We develop a demand model with switching costs, quality beliefs, and inattention, and estimate it using a field experiment. We find that (i) requiring active choice barely increases Bing’s market share; (ii) Google users paid to try Bing update positively about its quality and many prefer to continue using it; (iii) many Google users defaulted into Bing do not switch back, consistent with inattention. Counterfactuals suggest that eliminating demand frictions doubles Bing’s market share. Successful remedies expose users to alternative search engines, while data sharing mandates have small effects.

*We are grateful to Anlei Dong, Glenn Ellison, Chiara Farronato, Jingyi Guo, Paul Heidhues, Minha Hwang, Aadharsh Kannan, Widad Machmouchi, Markus Mobius, Sarah Moshary, Aviv Nevo, Michael Schwartz, Fiona Scott Morton, Steve Tadelis, Catherine Tucker, Mike Whinston, and audiences at Amazon, Charles River Associates, Chicago Booth, DMA and Beyond, NBER SI Digital Economics, Harvard-MIT IO Seminar, HBS Markets and Competition Conference, Northwestern Antitrust Conference, NYU Stern IO day, Paris School of Economics, Stanford SITE, Stanford Behavioral Seminar, University of Michigan, University of Zurich, Wharton, UVA, VIDE Seminar, Temple University, and TSE Economics of Platforms Seminar for helpful comments. We thank Chris Karr and Audacious Software for dedicated work on the Search Extension browser extension, and we thank Chiara Farronato and Andrey Fradkin for allowing Search Extension to use code developed for their work. We thank Shotaro Beppu, Jack Cenatempo, Juan Carlos Cisneros, Wonjoon Choi, Grace Coogan, Xingtong Jiang, Luca Moreno-Louzada, Sameer Nair-Desai, Shiqi Yang, and Wanxi Zhou for exceptional research assistance. We are grateful to the Sloan Foundation, the Sloan Research Fellowship, the Stanford Institute for Economic Policy Research (SIEPR), and the Business Economics and Public Policy Department at the Wharton School of the University of Pennsylvania for generous support. We also gratefully acknowledge Microsoft Bing for sharing data with us. The experiment was approved by the MIT Committee on the Use of Human Subjects (Protocol # 2308001088) and was registered in the American Economic Association Registry for randomized control trials under trial AEARCTR-0012884; the pre-analysis plan is available from <https://www.socialscisearchregistry.org/trials/12884>. Disclosures: Allcott, Castillo, and Musolff have previously worked at Microsoft Research. Gentzkow has done litigation consulting for Google and has been a member of the Toulouse Network for Information Technology, a research group funded by Microsoft. Allcott: Stanford University and NBER, email: allcott@stanford.edu. Castillo: University of Pennsylvania and NBER, email: jccast@upenn.edu. Gentzkow: Stanford University and NBER, email: gentzkow@stanford.edu. Musolff: University of Pennsylvania, email: lmusolff@wharton.upenn.edu. Salz: Massachusetts Institute of Technology and NBER, email: tsalz@mit.edu.

1 Introduction

Search engines are the gateway to the internet, the starting point for 69 percent of all online activities and 44 percent of online purchases.¹ According to a bipartisan report by the [US House Judiciary Subcommittee on Antitrust \(2020\)](#), search engines are part of “the infrastructure of the digital age” and have the potential to “pick winners and losers throughout [the] economy.” Due to this key position in the online ecosystem, better search engines can unlock substantial benefits for consumers and firms. Strengthening competition and improving efficiency in the web search market are therefore important policy goals.

Google holds approximately 90 percent of the global web search market ([StatCounter 2024b](#)). Antitrust authorities allege that Google has secured this dominance through anticompetitive practices—such as contracts that make it the default search engine—and by economies of scale in data.² Following this reasoning, the US Department of Justice sued Google in 2020 ([Department of Justice 2020](#)), and in 2024 the DC district court ruled that Google is a monopolist that has engaged in anticompetitive behavior ([DC District Court 2024](#)). Google contends, however, that its success reflects its high quality, that competition is “only a click away” ([Page 2012](#)), and that returns to scale in data are modest ([Varian 2015](#)).

We study two questions at the core of this debate. First, why is Google’s market share so large? It may be due simply to higher quality. It might stem from users’ lack of exposure to alternative search engines, which leads to misperceptions about their quality. It might also result from default effects, which could strengthen Google’s position directly through switching costs and inattention, and indirectly by limiting exposure to alternatives. Economies of scale in data could reinforce many of these advantages. Second, how would widely discussed policy interventions such as active choice screens, alternative defaults, and mandatory data sharing impact the market?

To answer these questions, we first develop a model of demand for search engines. Based on our model, we design and implement a field experiment that allows us to identify the sources of Google’s market power and estimate our model parameters. Using internal Bing search logs, we then estimate economies of scale in data. Finally, we simulate counterfactuals to evaluate the impacts of proposed policy interventions. Throughout our analysis, we focus on the US desktop search market and the competition between Google and Bing, which account for 93 percent of that market ([StatCounter 2024b](#)).

In our demand model, internet users make a binary choice between two search engines (Google and Bing) in each period. In an initial period, their web browser determines the initial default search engine. There are two sources of inertia in switching away from the default: a switching cost that reduces switching

¹See [SEO Statistics \(2024\)](#) and [Ecommercedb \(2024\)](#).

²The [UK Competition and Markets Authority \(2020\)](#) has expressed concerns that “we are currently in a catch 22 situation, whereby demand-side remedies would not be sufficiently effective until search engines have access to the level of search data needed to improve their results.” The EU’s competition authority brought a case against Google in 2003, which resulted in a \$4.1 billion fine for Google and the implementation of a choice screen ([European Commission 2018](#)).

by marginal users, and inattention that prevents some users from switching even if they would strongly prefer to. Users may also begin with incorrect beliefs about the quality of the other search engine. Users switch away from the default if they are paying attention and if the perceived quality improvement of the other search engine outweighs the switching cost.

We designed our experiment both to provide model-free evidence on the key forces and to identify the model's parameters. We recruited a sample of 2,354 desktop internet users from Prolific, a high-quality online survey panel, and asked them to take a baseline survey ("Survey 1") and a follow-up survey ("Survey 2") two weeks later. On Survey 1, participants answered questions about their web search preferences and installed a web browser extension. The extension recorded search engine queries and clicks from two weeks before the first survey until two months after. We also conducted an exit survey at the end of the study.

We randomized participants into a control group and three treatment groups. The Active Choice group was asked what search engine they would like to be their default and then received detailed guidance to implement their choice. The Default Change group was offered \$10 to change their default for two days and then guided similarly; afterwards they received no further instructions or incentives. The Switch Bonus group was offered a payment to change their default search engine for 14 days, guided to implement the change if they agreed to the offer, and then asked to make an active choice on Survey 2. The majority was offered a \$10 payment; smaller subsets were offered either \$1 or \$25. Google users in the \$10 Switch Bonus group were further randomized into two interventions implemented by our browser extension: (i) the Ranking Degradation condition, which decreased the relevance of Bing's results by reversing the order of organic search results on any result page, and (ii) the Ad Blocking condition, which removed most ads from Bing's search result pages.

Before the experiment, 96 percent of participants used Google for the majority of their searches. As expected, the Control intervention did not materially affect market shares.

The Active Choice intervention also had almost no effect on Google users: only 1.1 percent chose to switch to Bing. The small switching rate suggests that eliminating switching costs and inattention would not meaningfully reduce Google's market share. By contrast, among Bing users, 14 percent chose to switch to Google when required to make an active choice, suggesting that switching costs and inattention have a larger impact on Bing users.

In the \$10 Switch Bonus group, 58 percent of Google users switched to Bing in exchange for our payment. Exposure to Bing increased users' self-reported perceptions of its quality by 0.6 standard deviations. Of those who switched to Bing, 33 percent kept using Bing after making an unincentivized active choice on Survey 2. As these users only differ from Active Choice users in having had exposure to Bing, we therefore interpret this result as learning about Bing, either about its quality or about how to use it. Our exit survey confirms this: 64 percent of participants who actively decided to keep using Bing reported that it

was better than expected, and 59 percent reported that they had gotten accustomed to using Bing. These answers suggest that Google users' lack of experience with Bing is a significant driver of Google's large share at baseline in our sample. For Bing users, the discrepancy between choices before and after exposure to Google is negligible, indicating that Bing users are generally well-informed about Google's quality. The survey results again support this interpretation: updating about Google is less pronounced and mostly statistically insignificant.

The Default Change intervention increased Bing's market share among baseline Google users, and this effect persisted until the end of our experiment. Of the users in this group, 81 percent accepted the \$10 to switch to Bing for two days. Among compliers, we observe a gradual decline in Bing's market share from the initial 100 percent, to 66 percent after one week, 60 percent after two weeks, and 46 percent after two months. Our model suggests two reasons for the continued use of Bing by these users. First, like Switch Bonus group participants, their valuation of Bing increases due to experience. Second, some participants may continue to prefer Google but not switch back due to persistent inattention. Our exit survey confirms that both play a role: 35 percent of users who kept using Bing report doing so because they prefer it, while 44 percent report that they forgot to switch back or were too lazy to do so. This result has two implications. First, defaults create a lasting mismatch between preferences and choices. Second, changing defaults can induce learning about unobserved product quality, leading to lasting effects by altering users' perceptions.

Our price treatments uncover substantial heterogeneity in participants' willingness to accept switching search engines. Among Google users, a \$1 payment to switch to Bing for two weeks raised Bing's market share to 32 percent, meaning many users are close to indifferent. With a \$10 payment, Bing's market share increased to 64 percent. At \$25, Bing's market share only increased to 74 percent, meaning many users have strong preferences for Google.

The results from our Ranking Degradation intervention, although somewhat noisy, suggest at most a moderate demand response to search result relevance. While this intervention substantially affected perceptions—it significantly reduced the reported quality of search result relevance and overall search engine quality—we do not detect statistically significant changes in market shares. According to our point estimates, Ranking Degradation reduced Bing's market share by 3.4 percentage points (standard error=2.9). Internal experiments at Google similarly suggest that "a significant quality depreciation by Google would not result in a significant loss of revenues" ([DC District Court 2024](#)). However, given our standard errors, our results are consistent with a reduction of up to 9 percentage points.

We use the data from our experiment to estimate our model by generalized method of moments. We find that switching costs are negligible, but 34 percent of users are persistently inattentive. In the model, the median Chrome user would have to be paid \$3.20 to use Bing instead of Google for two weeks. After two weeks of experience with Bing, the required payment shrinks to \$2.99. Although this dollar difference

is small, many users have weak preferences over search engines, so learning shifts market shares significantly. The parameters measuring responsiveness to search result quality and ad-loads are modest in size and statistically indistinguishable from zero.

We complement our demand side with estimates of returns to scale in data using internal search logs from Bing. Specifically, we estimate how click-through rates (a standard measure of result relevance) improve for unseen novel queries as Bing serves more results and collects more data. Returns to data are positive but diminishing, with a logarithmic relationship between cumulative queries and the click-through rate. This relationship predicts that if Bing had access to Google’s data, click-through rates would increase from 23.5 percent to 24.8 percent. While there are many potential sources of scale economies in web search, such as in web indexing and other fixed costs, this analysis specifically isolates benefits relevant for proposed antitrust remedies: those from more click-and-query data. However, this analysis requires strong assumptions, and is therefore more speculative than the rest of our findings.

We consider several counterfactuals. First, we simulate the effect of shutting down all demand-side frictions—switching costs, inattention, and misperceptions about quality. Bing’s market share increases from 11 to 24 percent, showing that these frictions create a substantial barrier preventing Bing from increasing its market share. Consumer surplus increases by \$6 per consumer per year.

Next, we simulate the effects of active choice screens on Chrome that appear when the web browser is first installed, as required in the European Union. To this end, for Chrome users only, we shut down switching costs and inattention, but still allow users to misperceive quality. Driven by limited effects of our Active Choice intervention, our model predicts that choice screens would increase Bing’s market share by only 1.3 percentage points. This small increase underscores that, while choice screens can address certain frictions, they only have a limited impact because they do not eliminate the larger barrier to competition: users’ misperceptions of Bing’s quality.

These results suggest that if regulators want to significantly impact market shares, they should recognize search engines as experience goods (Nelson 1970) and consider how interventions impact consumers’ exposure to Bing. Thus, we next use our model to measure the effects of an intervention that could increase exposure. If Google were prevented from being the default search engine, Bing could become the default on all browsers, increasing Bing’s market share by 39 percentage points. However, consumer surplus would decrease by \$73 per consumer per year, because a large number of users would use Bing despite strongly preferring Google.

Our results so far highlight a conundrum for competition policy. While choice screens increase consumer surplus by a modest amount, they have almost no effect on market shares. Changing defaults, on the other hand, has a large effect on market shares, but only at the expense of a large decrease in consumer surplus. This raises the question of whether a policy exists that can reduce market concentration without lowering

consumer surplus. One possible approach is to mandate that a non-dominant firm—in this case, Bing—be set as the default on all browsers upon installation, followed by a requirement that browsers present a choice screen *after some time*. This would allow users to experience Bing before making an active choice. Such a policy would reduce Google’s market share by 15 percentage points, while leaving consumer surplus essentially unchanged. Thus, a delayed choice screen could avoid the harm caused by simply setting Bing as the default while reaping the potential benefits from a less concentrated market, such as increased investment incentives and fewer harms on the advertising side (which we do not model explicitly).

Finally, in a more speculative analysis, we account for the feedback effects from endogenous result quality and we simulate the effects of providing Google’s search results and click data (“click-and-query” data) to Bing, using the economies of scale estimated from the Bing search logs. Neither feedback effects nor data sharing substantively affect market shares or consumer surplus. This follows from two earlier results: data sharing has only a small effect on Bing’s result quality in the Bing search data, and Bing’s result quality has a small effect on market shares in the experiment.

What do our results imply for the discussion surrounding Google’s dominance? First, we find that most people *do* prefer Google over Bing, and some strongly. At the same time, Google significantly benefits from frictions that raise its market share beyond the efficient level. The debate has largely focused on the \$26 billion that Google spends annually to secure its default position on browsers and Android devices ([DC District Court 2024](#)), and our findings confirm that defaults play an important role. However, our results suggest that the power of Google’s default position on Chrome does not stem from directly preventing users from choosing Bing, since most consumers at least think they prefer Google. Instead, Google’s default position is effective because it ensures that users are never exposed to Bing, and hence never learn about it. Such learning would permanently lower Google’s market share in our model. Our findings suggest that regulators and antitrust authorities can increase market efficiency by considering search engines as experience goods and designing remedies that induce learning. This conclusion may be of broader relevance. Prior theoretical literature has shown that incumbents might benefit from favorable user expectations ([Schmalensee 1982](#)). Our results imply that this theory is indeed relevant in this market.

Our results have several important limitations. First, we focus only on desktop search in browsers because parts of our experiment cannot be implemented on mobile. Desktop is important *per se*, representing 55 percent of web traffic in 2023 ([StatCounter 2024a](#)), but switching costs could be higher on mobile and on non-browser search bars integrated into Windows, Android, and iOS. Second, our experiment sample is more educated, has higher income, and is more white than the population of US adults, and it may not be representative on unobserved factors such as price sensitivity or computer literacy. Third, our economies of scale analysis requires strong identifying assumptions.

Our work contributes to several related literatures. First, we contribute to the literature on competition

and antitrust in web search (Ostrovsky 2021; Vázquez Duque 2022; Decarolis, Li, and Paternolli Forthcoming; Hovenkamp 2024; Bhargava, Kraemer, and Wipusanawan 2025) and the surrounding policy discussion (Patterson 2013; Stigler Committee on Digital Platforms 2019; UK Competition and Markets Authority 2020; Dinielli et al. 2023; Heidhues et al. 2023).³ Two papers are particularly related to our work. Decarolis, Li, and Paternolli (Forthcoming) observationally investigates antitrust remedies imposed by European and Russian regulators. They find small effects from the EU choice screen, consistent with our Active Choice results. Vázquez Duque (2022) conducts a survey experiment on Amazon Mechanical Turk in which participants choose search engines either in an active choice or a default treatment. He finds that active choice has a small effect on market shares, and that misperceptions significantly contribute to Google’s high market share. Our work extends Vázquez Duque (2022) in three main ways: we conduct a field experiment based on incentivized real-world choices, our browser extension allows us to implement additional treatments essential for disentangling the sources of Google’s market power, and we model counterfactuals that speak directly to policy.

Second, we extend work on the competitive effect of choice frictions, including switching costs (Klemperer 1987; Farrell and Klemperer 2007) and defaults in the presence of inattention (DellaVigna and Malmendier 2006; Carroll et al. 2009; Handel 2013; Ericson 2014; Ho, Hogan, and Scott Morton 2017; Andersen et al. 2020; Fowlie et al. 2021; Einav, Klopach, and Mahoney Forthcoming; Miller, Sahni, and Strulov-Shlain 2023; Brot-Goldberg et al. 2023; Lee and Musolff 2023). Our results highlight an important new role of defaults in a market setting. Even when switching costs and inattention are relatively small, defaults matter by preventing consumers from gaining experience with an alternative whose quality they initially underestimate. Like Agte et al. (2024), we thus find that switching costs and misspecified beliefs interact.

Third, we extend the empirical literature that studies experience goods (Erdem and Keane 1996; Akerberg 2003; Israel 2005; Crawford and Shum 2005; Dickstein 2021) by showing that overly pessimistic consumer beliefs about the quality of rivals help entrench dominant firms.

Fourth, we extend previous empirical work on economies of scale in search (Chiou and Tucker 2021; He et al. 2017; Azevedo et al. 2020; Schaefer and Sapi 2023; Klein et al. 2023), and in data more broadly (Bajari et al. 2019; Tucker 2019), by combining our estimates of economies of scale with experimental estimates to quantify the equilibrium implications of antitrust remedies and the resulting welfare effects for consumers.

Fifth, we contribute to a literature that experimentally studies digital markets, including studies on consumer surplus from social media (Allcott et al. 2020), price salience (Blake et al. 2021), addiction to digital services (Allcott, Gentzkow, and Song 2022), substitution pattern across online services (Aridor Forthcoming).

³Most existing work on the search engine market has focused on advertising (Varian 2007; Edelman, Ostrovsky, and Schwarz 2007; Athey and Ellison 2011; Blake, Nosko, and Tadelis 2015). In addition, there are studies assessing the value of digital services that don’t charge prices to consumers (Brynjolfsson, Collis, and Eggers 2019), of which search engines are an important example.

ing), advertising (Goli et al. 2025; Brynjolfsson et al. Forthcoming; Wernerfelt et al. 2024; Katz and Allcott 2025), and the welfare consequences of platforms when people experience fear of missing out (Bursztyn et al. 2023). To facilitate such studies, Farronato, Fradkin, and Karr (2024) introduced an open source browser extension, whose code helped us develop the extension for this project.

2 Model

We now present our model of search engine demand, which guides our experimental design.

2.1 Search engine choices

Consumers indexed by i make a binary choice between two search engines $j \in \{B, G\}$ in two-week periods indexed by t . We think of this choice as determining the search engine for both direct navigation and address bar searches.⁴ Consumer i 's search engine choice in period t is denoted by y_{it} .

Each consumer has an exogenously set web browser (Chrome or Microsoft Edge), which determines her default search engine, i.e., the search engine used for address bar searches when a browser is first installed. We denote the default search engine by d , and the alternative search engine by $-d$. For Chrome users, the default is Google ($d = G$), and the alternative search engine is Bing ($-d = B$). For Microsoft Edge users, the default is Bing and the alternative search engine is Google.

The payment agent i receives for using search engine j in period t is p_{ijt} . In real life, $p_{ijt} = 0$, but our experimental interventions will vary prices. Variables a_j^* , r_j^* , and ξ_j^* refer to j 's ad load, search result relevance, and other unobserved characteristics respectively. We define $\zeta_j^* := \alpha a_j^* + \rho r_j^* + \xi_j^*$ as j 's "quality." Stars denote the fact that these are true quantities—below, we allow users to be imperfectly informed about them.

Each period's flow utility from j is

$$u_{ijt}^* = \eta p_{ijt} + \zeta_j^* + \chi_{ij}, \quad (1)$$

where χ_{ij} is an idiosyncratic preference shifter that does not vary across time periods.

Users may be imperfectly informed about search engine quality. We use $\mathbb{E}_{it}$ to denote agent i 's expectation of different quantities at time t . Thus, the quality agent i perceives at time t is given by $\mathbb{E}_{it}[\zeta_j] := \alpha \mathbb{E}_{it}[a_j] + \rho \mathbb{E}_{it}[r_j] + \mathbb{E}_{it}[\xi_j]$. Although the true quality ζ_j^* is constant, the perceived quality $\mathbb{E}_{it}[\zeta_j]$ depends

⁴In our experiment, only 6.6 percent of users perform more than 10 percent of searches on a non-default search engine.

on time because users' perceptions may change, as we explain below. The perceived flow utility is given by

$$u_{ijt} = \eta p_{ijt} + \mathbb{E}_{it}[\zeta_j] + \chi_{ij}. \quad (2)$$

We assume that users have correct beliefs about their default search engine d , so $\mathbb{E}_{it}[\zeta_d] = \zeta_d^*$ and $u_{idt} = u_{idt}^*$. Users that have never chosen the alternative search engine $-d$, by contrast, may be imperfectly informed about its quality. In that case, perceived quality $\mathbb{E}_{it}[\zeta_{-d}]$ takes a different value $\tilde{\zeta}_{-d} := \alpha \tilde{a}_{-d} + \rho \tilde{r}_{-d} + \tilde{\xi}_{-d}$. We assume that this value does not vary across imperfectly informed users. After one period of experience with $-d$, consumers become fully informed, so their perceived quality becomes $\mathbb{E}_{it}[\zeta_{-d}] = \zeta_{-d}^*$.

Defaults could matter because there are two sources of inertia. First, consumers face a switching cost σ for getting to the choice screen and changing the default. Second, consumers may be inattentive. A fraction ϕ of users (called “permanently inattentive”) never pay attention, so they always stay with their default. The remaining fraction $1 - \phi$ of users are probabilistically inattentive. In each period, with exogenous probability π they are attentive and consider the choice between search engines, and with probability $1 - \pi$ they are inattentive.

In period t , if inattentive, consumers stick with the search engine they used in the previous period: $y_{it} = y_{i,t-1}$. If attentive, the consumer chooses the search engine that maximizes utility over an infinite horizon with per-period discount factor δ :

$$y_{it} = \arg \max_{j \in \{G, B\}} \{u_{ijt} - \sigma \mathbf{1}_{j \neq y_{i,t-1}} + \delta V_{i,t+1}(j)\}. \quad (3)$$

where $V_{i,t+1}(j)$ is the perceived continuation value after having chosen search engine j .

We make three assumptions that simplify this dynamic switching problem into an effectively static decision. First, consumers do not perceive any uncertainty about quality ζ_{-d} , so there is no option value to exploration. Concretely, consumers make decisions based on beliefs that are degenerate at $\mathbb{E}_{it}[\zeta_d]$. Second, as this is consistent with our experimental estimates, we assume consumers weakly underestimate the quality of the alternative search engine ($\tilde{\zeta}_{-d} \leq \zeta_{-d}^*$), so experience with $-d$ weakly increases its market share. Finally, we assume that by the start of the experiment all participants that are not permanently inattentive have made an attentive choice, so that market shares are in steady state at $t = 0$. This approximates a world in which the time between browser installation and the beginning of the experiment is long.⁵

Given these assumptions, and since idiosyncratic preferences remain constant over time, the consumer's decision is effectively static: if it is optimal to switch in the future, it is optimal to switch immediately. Thus,

⁵We believe that the first assumption is psychologically realistic, and also consistent with what we observe in the Active Choice treatment. The second assumption is also consistent with our experimental estimates of experience effects. The final assumption is consistent with the lack of market share trends before our experiment and in our Control group.

attentive consumers permanently choose either B or G , where they account for the perceived discounted utility $u_{ijt}/(1 - \delta)$ of each search engine. Equation (3) thus simplifies to

$$y_{it} = \arg \max_{j \in \{G, B\}} \left\{ \frac{u_{ijt}}{1 - \delta} - \sigma \mathbf{1}_{j \neq y_{i,t-1}} \right\}. \quad (4)$$

At baseline—before the experiment starts—prices are equal to zero.⁶ Hence, the perceived discounted utility from permanently continuing with the default search engine and from permanently switching to the alternative search engine are, respectively,

$$\frac{\zeta_d^* + \chi_{id}}{1 - \delta} \quad \text{and} \quad \frac{\tilde{\zeta}_{-d} + \chi_{i,-d}}{1 - \delta} - \sigma. \quad (5)$$

We now define variables for differences between the alternative search engine and the default search engine: $\Delta v_t := v_{-d} - v_d$ for any variable v . Concretely, $\Delta p_{it} := p_{i,-d,t} - p_{idt}$, $\Delta \zeta^* := \zeta_{-d}^* - \zeta_d^*$, and $\Delta \chi_i := \chi_{i,-d} - \chi_{i,d}$. We also define $\Delta \tilde{\zeta} := \tilde{\zeta}_{-d} - \zeta_d^*$, since users always perceive the utility of the default search engine correctly. With this notation and after differencing the expressions in equation 5, an attentive consumer switches to $-d$ if

$$\Delta \tilde{\zeta} + \Delta \chi_i - \sigma(1 - \delta) > 0. \quad (6)$$

Therefore, the probability that a consumer that pays attention chooses $-d$ is $\mathbb{P}(\Delta \tilde{\zeta} - \sigma(1 - \delta) > -\Delta \chi_i)$. Defining $S(\cdot)$ as the cumulative density function of $-\Delta \chi_i$, this becomes $S(\Delta \tilde{\zeta} - \sigma(1 - \delta))$.

A fraction ϕ of users is permanently inattentive, and thus will keep d forever. The remaining fraction $(1 - \phi)$ will eventually pay attention at some point and choose the search engine that maximizes perceived utility. Thus, at the beginning of our experiment ($t = 0$), $-d$'s market share is

$$s_{-d,0} = (1 - \phi)S(\Delta \tilde{\zeta} - \sigma(1 - \delta)). \quad (7)$$

Equation (7) illustrates four reasons why Google might have high steady-state market share: (i) Google has a higher true quality than Bing ($\Delta \zeta^* < 0$), (ii) consumers perceive Bing to be worse than it actually is ($\tilde{\zeta}_B < \zeta_B^*$), (iii) Google is many users' initial default and the switching cost σ is large, or (iv) Google is many users' initial default and the fraction ϕ of permanently inattentive users is large.

Optimal choices are attained in a counterfactual where consumers learn the true ζ and make active choices with zero switching cost. In that scenario, the market share of $-d$ is $s_{-d,0} = S(\Delta \zeta^*)$.

⁶In practice, users can earn rewards for Bing searches and redeem them for Microsoft products. However, as rewards are modest, we do not model them directly. This modeling decision effectively means that rewards are captured by the quality term ζ .

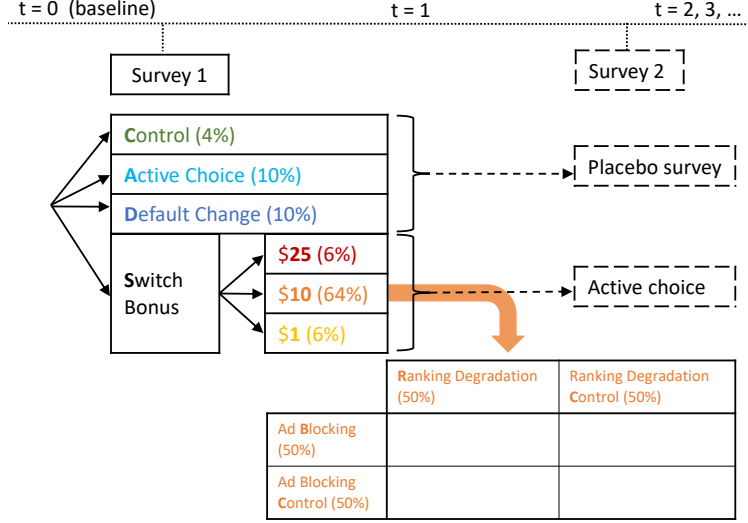


Figure 1: **Experimental Design**

Notes: This figure illustrates our experimental design. The Control group was guided through how to change their bookmarks. The Default Change group was paid \$10 to switch their default from Google to Bing for two days. The Active Choice group was asked about their preferred default and then guided to implement this choice. The Switch Bonus treatment group was asked to switch search engines in return for a bonus payment of either \$1, \$10, or \$25. The \$10-group was further randomized in a two-by-two factorial design, which varied whether ads were blocked and whether search results are degraded. The Switch Bonus group made an active choice after the 14-day incentive period.

3 Experimental Design

We now present a field experiment that we designed to measure the parameters of our demand model. In Section 5 we show formally how our experimental treatments identify each parameter.

3.1 Overview

Figure 1 illustrates the experimental design and timeline. There are two surveys. Survey 1 took place immediately after recruiting. The invitation to Survey 2 was sent the morning of the 15th day after participants complete Survey 1. The experiment ended two months after Survey 1. We sent an exit survey to some users at the end of the experiment.

Recruitment, screening, and demographics. From March 19th to April 2nd, 2024, we recruited participants from the Prolific online platform, enforcing balance by gender. To qualify for the study, participants had to be US residents and at least 18 years old.

Survey 1 began with screening questions concerning the participant’s device, web browser, and search engine use. To obtain a survey-based measure of their current address bar search engine, we asked participants to search for the term “potato” through their browser’s address bar and report the search engine that

they were directed to. Unless users intentionally changed the address bar search engine, it is the browser’s default search engine. The survey then separately asked, “What search engine do you usually use on this web browser?” in case users do not usually search via the address bar.

Participants could continue with Survey 1 only if they accessed the survey through a desktop or laptop computer using Edge or Chrome and they reported that (i) on their current device, they exclusively use the current web browser, (ii) they do not frequently share that computer with other people, (iii) their address bar search engine is either Google or Bing, and (iv) the search engine they usually use is either Google or Bing. We restrict the analysis to participants with a “consistent baseline search engine”: the search engine they report they usually use is also (i) the search engine they reported as the address bar search engine (ii) the search engine for more than half of searches recorded by our browser extension before installation, and (iii) the address bar search engine recorded by our browser extension. We also drop users with fewer than 10 recorded searches in the 20 days before installing our browser extension during Survey 1.

Participants who passed the screening questions and consented to participate were then asked demographic questions. This was followed by a series of questions eliciting opinions about Google and Bing, including why they use the search engine they usually use.

Search engine rating questions. For all participants on Survey 1 and some participants on Survey 2, we asked several *search engine rating questions*. We first asked people to rate Google versus Bing in terms of overall quality and on six specific dimensions: (i) relevance and ordering of search result links, (ii) features on search result pages (e.g., weather info), (iii) relevance of ads, (iv) AI chat, (v) privacy, and (vi) rewards or loyalty points. Possible answers were *Google is a lot better*, *Google is a little better*, *they are about the same*, *Bing is a little better*, and *Bing is a lot better*. The answer order was randomly flipped. We also included an attention check, to “please choose ‘Bing is a lot better’ if you are still paying attention.” The questions were presented in random order.

Search Extension. Participants were then asked to install Search Extension, a browser extension developed for this study. As a standard Chrome/Edge browser extension, it is unobtrusive and not visible on the browser interface after the installation. Search Extension records the dates, times, and information identifying the source (the address bar or the search engine website) of all searches on all general web search engines (google.com, bing.com, etc.) that take place after installation. Using the browser’s recorded search history, it also collects the same information for all searches made in the 20 days before installation.⁷ Additionally, for searches made after installation, the extension records whether the user clicked on an ad or an organic search result, and if so, the rank of the result.

⁷Throughout our analysis, we only use data for the two weeks prior to installation to harmonize the data with our experiment.

Search Extension includes two intervention functionalities that we turned on or off in treatment conditions described below. First, the Ranking Degradation functionality reverses the order of organic results on search result pages. Thus, on each page, the bottom results are moved to the top, and the top results are moved to the bottom. Second, the Ad Blocking functionality removes all ads that it detects on search result pages. Search Extension does not make users aware of these functionalities or whether they are turned on. These interventions occur in a split-second when the page loads and are imperceptible to the user.

Compensation. In addition to the incentive payments associated with each treatment, participants were paid a base payment of \$25: \$5 each for completing Survey 1 and Survey 2, \$5 for installing Search Extension, and \$10 for keeping Search Extension installed for two months after completing Survey 1.

3.2 Treatment Groups

Users were randomized into four treatment groups, one of which has further sub-treatments. Participants whose baseline default search engine was Google were randomized into all groups, with the proportions in Figure 1 and below. Since there are relatively few baseline Bing users, they were randomized into only two groups (A and S10CC, with equal probability, see below) to increase power. We did not include a Control group for Bing users because market shares are stable over time without intervention.

We now describe each group’s experience after installing Search Extension on Survey 1.

Control (group “C,” 4 percent of baseline Google users). As a placebo intervention, the Control group was shown information about how to change the bookmarks on their web browser, in a similar format to the treatment information the other groups receive. We correctly anticipated that this placebo intervention would not change search engine market shares.

Active Choice (group “A,” 10 percent of baseline Google users). The Active Choice group was told that we would now show them how to change the default search engine. To avoid experimenter demand effects, the survey clearly stated that “whether you change it or not is up to you.” The survey then asked, “when we get to the screen where you can set your default search engine, what would you like your default to be?” We call this the person’s desired default. The survey then showed people how to change the default search engine, asked people to copy and paste the settings page URL (to confirm that they were on the correct page), asked people to set the address bar default search engine to their desired default, and asked to confirm that they had done so or explain why not.

Default Change (group “D,” 10 percent of baseline Google users). The Default Change group was offered \$10 to switch their default to Bing and conduct at least 90 percent of their searches (minimum 4 searches in total) on Bing over the next 2 days. Participants were told we would show them how to change their default search engine, then asked: “Would you like to accept the additional \$10 to make Bing your primary search engine for the next 2 days?” As in Active Choice, the survey then explained how to change the default and asked people to copy and paste the settings page URL. Those accepting the offer were instructed to change the address bar default and to either confirm that they had done so or explain why not.

For these first three groups (C, A, and D), we wanted a second survey only to make sure that all groups have the same number of surveys. Thus, Survey 2 for those three groups is a “placebo survey” that simply reminds participants of the payments they get if they stay until the end of the experiment.

Switch Bonus (group “S,” 72 percent of baseline Google users). Following the same steps as for the Default Change group, the Switch Bonus group was offered payment to switch search engines for fourteen days and to make at least 90 percent of their searches (minimum 20 searches total) on the alternative search engine during this time. Payments p varied between \$1, \$10, or \$25, with 6, 64, and 6 percent probability, respectively. To avoid considerations of future inertia, participants were informed: “on the second survey in 14 days, we will remind you how to switch your default search engine.” On Survey 2, they were asked the search engine rating questions and then received the active choice intervention described above.

For baseline Google users, the \$10 Switch Bonus group (with 64 percent of the sample) was further factorized into a two-by-two matrix of two interventions implemented by Search Extension. First, **Ranking Degradation** (group “R,” 50 percent of S10 group) users experienced the ranking degradation functionality on Bing. Second, **Ad Blocking** (group “A,” 50 percent of S10 group) users experienced the ad blocking functionality on Bing. We refer to the 25 percent subset of the S group assigned to the Ranking Degradation Control and Ad Blocking Control groups as group S10CC. We made the \$10 Switch Bonus group relatively large and the Ad Blocking and Ranking Degradation conditions relatively forceful because we expected limited power to detect the effects of these two interventions on market shares.

3.3 Exit Survey

At the end of the experiment (two months after Survey 1), we sent an exit survey to a random subset of participants in the Switch Bonus and Default Change groups whose original search engine was Google and who kept Bing after the incentive period. Eligible participants were offered \$5 for completing it. In a free-form text response field, both groups were asked “Over the past 6 weeks, our records show that you have continued to use Bing on this browser. Why?” Subsequently, they were asked “Why did you decide to keep using Bing? Please choose all that apply.” The available options for this question were different for both

groups. For the Switch Bonus group, the possible answers were (i) *Before the experiment, I had always wanted to use Bing but hadn't gotten around to it*, (ii) *Bing was better than I thought it would be*, (iii) *I got accustomed to Bing*, and (iv) *Other*, with forced free form response. For the Default Change group, the possible answers were (i) *I wanted to keep using Bing*, (ii) *I forgot to change back to Google*, (iii) *Changing back to Google was too much effort*, and (iv) *Other*, with forced free form response.

3.4 Pre-Analysis Plan

We submitted our final pre-analysis plan (PAP) to the AEA RCT registry in February 2024, the month before data collection began.⁸ The PAP specified how we would construct the basic experimental results, by presenting mockups of Tables 1, 2, 4, A1, and A2, and Figures 1, A2, 2, 3, and 4. Our tables and figures follow the PAP's mockups except for non-substantive reformatting, the addition of one row in Table 1, presenting a subset of the results in Figure 2 (Appendix Figure A1 presents the full pre-registered figure), and a sample restriction in Table 4 and Panel A of Figure 4.⁹

4 Experimental Results

4.1 Data

Table 1 shows the sample sizes at each point of the experiment. The sample reported in each row is a strict subset of the row above. The final row reports 1,461 people who kept Search Extension installed until the end of the experiment (two months after Survey 1). These participants form our final analysis sample for all results tables and figures.

Table 2 presents summary statistics. Panel A shows demographic covariates for the sample and for US adults. Our sample is nearly balanced on gender but younger, more educated, more white, and higher-earning than the US adult population. It also includes a higher share of desktop Google users. Panel B shows that participants conduct at least one search on over 60 percent of days, averaging about 11 searches daily. Using pre-experiment browser history data collected by Search Extension, we compare search behavior before and during the experiment (i.e., the eight weeks after Survey 1). We find no statistically significant difference in either the frequency of search days or the average number of daily searches between these periods, suggesting that any substitution to other browsers or devices is limited (Panel B).¹⁰ Panel C shows

⁸See <https://www.socialscisceregistry.org/trials/12884>.

⁹That sample restriction reflects a shift from intent-to-treat (ITT) effects in the PAP to treatment-on-the-treated (TOT) effects in Table 4 and Panel A of Figure 4. Appendix A.1 presents the pre-registered ITT versions.

¹⁰We further investigate the number of daily searches by treatment group in Appendix Table A4 and Appendix Figure A4. The only statistically significant result we find is that those baseline Google users that we are paying \$10 to use Bing search more during the incentive period, perhaps driven by the need to conduct more searches to find the right result. If users had substituted away to other devices, we would have found the opposite.

Table 1: Sample Sizes

	Sample size
US Prolific users not in pilots	45,219
Saw study advertisement	5,280
Started Survey 1	4,217
Passed screening questions	2,737
Consented	2,736
Finished Survey 1	2,518
Not rejected for multiple responses	2,390
Installed Search Extension	2,354
At least 10 baseline searches	1,914
Consistent baseline search engine	1,726
Finished Survey 2	1,660
Kept Search Extension 2 weeks after Survey 2	1,577
Kept Search Extension 2 months after Survey 1	1,461

Notes: This table presents sample sizes at each stage of the experiment. “Not rejected for multiple responses” drops users who tried to retake the survey after being disqualified in their first attempt. “Consistent baseline search engine” requires that the same search engine (Bing or Google) is (i) reported on Survey 1 as the search engine used for address bar search, (ii) reported as the search engine they usually use, (iii) the search engine for more than half of searches recorded by Search Extension before installation, and (iv) the address bar default recorded by Search Extension before installation. The sample in each row is a strict subset of the row above.

that after limiting to users with a consistent baseline search engine, there is only a minimal amount of multi-homing in our sample. Panel D cross-tabulates users by browser and by pre-experimental search engine.

In Appendix A, we present tests of balance and differential attrition. Treatment assignment is statistically balanced on observables; see Appendix Table A1. For baseline Bing users, completion rates are balanced across treatment groups. For baseline Google users, however, completion rates are statistically different: the Active Choice group has a slightly higher completion rate, while the Control group has a slightly lower completion rate. These differential completion rates cannot be explained by differential financial incentives or burdens to participation, so we believe them to be idiosyncratic. Differential attrition could affect our market share results only if attriters had different likelihoods of switching between Google and Bing. However, Appendix Table A3 shows that attriters versus stayers did not differ in their baseline ratings of Google versus Bing.

The attention check embedded in the search engine rating questions is passed by 96 percent of Google users and 100 percent of Bing users.

4.2 Initial Survey Ratings of Google and Bing

Figure 2 shows how participants rate search engine quality on Survey 1. Google users strongly prefer their search engine overall and across almost all specific dimensions of quality. The only exception is rewards, for which they express a slight preference for Bing—consistent with the fact that using Bing can earn

Table 2: Summary Statistics

Panel A: Sample Demographics		
	<i>Analysis sample</i>	<i>U.S. adults</i>
Income (\$000s)	56.45	40.86
College	0.58	0.33
Male	0.45	0.49
Age	36.39	48.16
White	0.60	0.32
Baseline Google user	0.96	0.82
Panel B: Search Activity		
	<i>Before Experiment</i>	<i>During Experiment</i>
	Mean (SE)	Mean (SE)
Fraction (%) of days with positive search	62.00 (3.50)	65.22 (3.29)
Daily searches	11.11 (1.78)	10.91 (1.16)
Panel C: Pre-Experiment Search Engine Share		
	<i>Share (%) of Google Searches</i>	
Google users	97.85 (0.14)	
Bing users	3.47 (0.84)	
Panel D: Users by Browser and by Pre-Experiment Search Engine (percentages)		
	<i>Google</i>	<i>Bing</i>
Chrome	94.32	1.23
Edge	1.16	3.29
Total	95.55	4.45

Notes: This table provides summary statistics of participants' search behavior both before (pre-Survey 1) and during the experiment (between Survey 1 and eight weeks after Survey 1). The statistics show the mean and standard error of individual-level averages.

users Microsoft Rewards worth up to \$10 per month, whereas Google offers no such rewards program. Meanwhile, Bing users only slightly prefer Bing on overall quality, result page features, AI chat integration, and privacy, but strongly prefer Bing in terms of rewards. Bing users rate both search engines similarly on relevance of links and ads.¹¹

4.3 Effects of Main Treatments

Figure 3 presents market shares by treatment group over time. To compute the market share at time t , we first compute the market share for each participant separately and then average shares across participants.¹²

¹¹We also asked participants why they chose Google or Bing. A significant share (53 percent of Bing users and 56 percent of Google users) attribute their usage to the browser's default. Bing users are less likely to report results page features (7 and 17 percent respectively) and relevance of links (26 and 69 percent respectively) as reasons for choosing their search engine. By contrast, Bing users are more likely than Google users to report AI integration (29 and 8.7 percent respectively) and rewards (66 and 1.9 percent respectively). Appendix Figure A2 presents full results.

¹²We obtain similar results if we define market shares by treating search engine choice as binary (defined as the search engine that the person used the most in period t and then taking the average of those binary choices). See Appendix Figure A6.

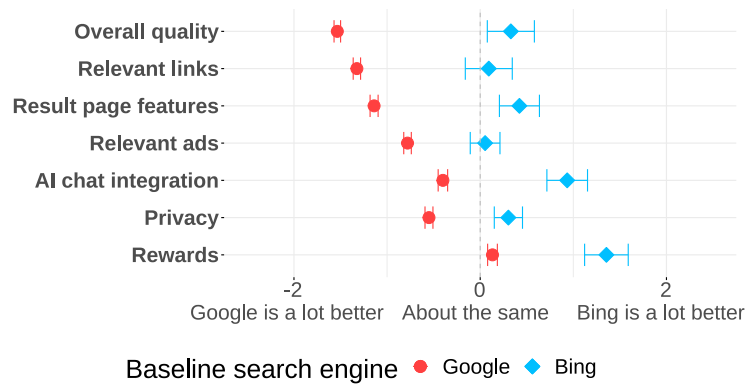


Figure 2: Initial Ratings of Google and Bing

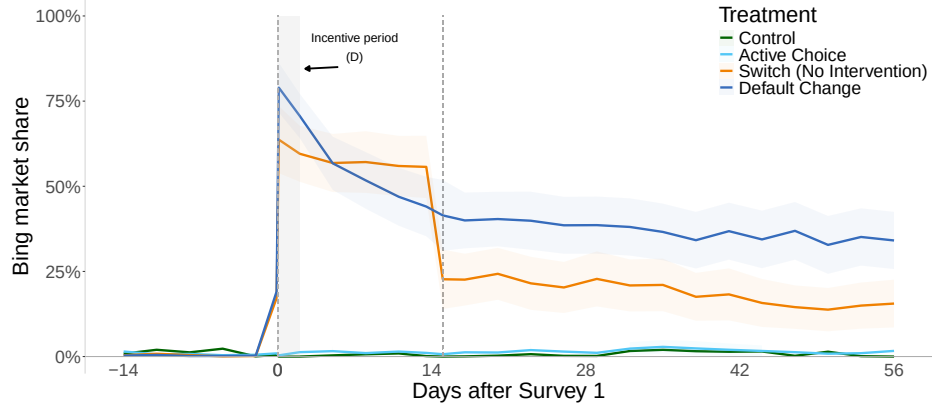
Notes: This figure presents average responses for baseline Google and Bing users in response to the following questions: “Overall, how would you rate the quality of Google relative to Bing?” and “How would you rate the quality of Google relative to Bing on the following dimensions?” Response options were “Bing is a lot better,” “Bing is a little better,” “They are about the same,” “Google is a little better,” and “Google is a lot better,” coded as 2, 1, 0, -1, and -2, respectively. Whiskers indicate 95 percent confidence intervals.

Panels (a) and (b) present results for baseline Google users. The \$10 Switch Bonus group presented here is limited to S10CC (“No Intervention”) participants, who did not experience the Ad Blocking or Ranking Degradation interventions. Panel (c) limits to baseline Bing users. It only includes the two groups to which they were assigned, Active Choice and Switch Bonus with no search extension intervention (S10CC).

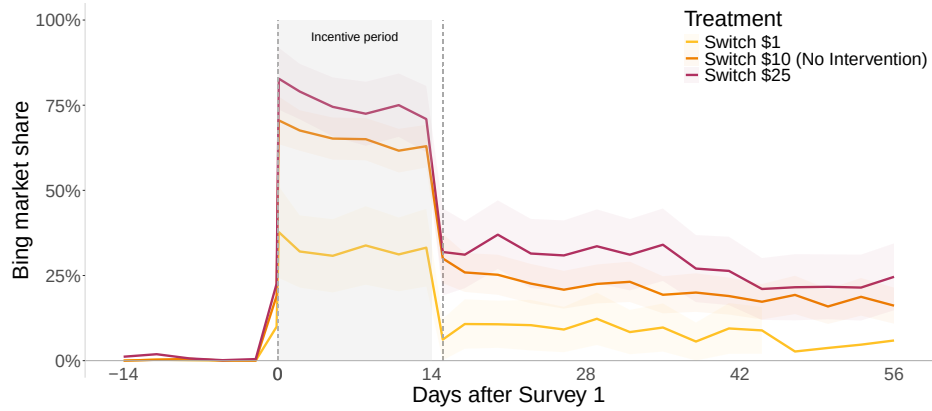
Baseline and Control. Panel (a) shows that the Bing market share among baseline Google users was virtually zero before the experiment, reflecting the fact that most of them exclusively used Google. We also see that control group users did not change their behavior during the experiment: 1 percent of searches after Survey 1 were made on Bing. Figure 3c shows that baseline Bing users did use Google occasionally before the experiment, but only for about 5 percent of searches.

Active Choice. Panel (a) also shows that the Active Choice intervention had almost no effect on the search engine usage of baseline Google users, increasing Bing’s market share from 0.7 percent to just 1.9 percent. This result is consistent with the relatively small effect of the choice screen that Google implemented on Android devices in the European Union (Decarolis, Li, and Paternollo Forthcoming). This means that most baseline Google users choose Google even when they are attentive and there are no switching costs, indicating that removing these frictions alone would not change Google’s large market share.

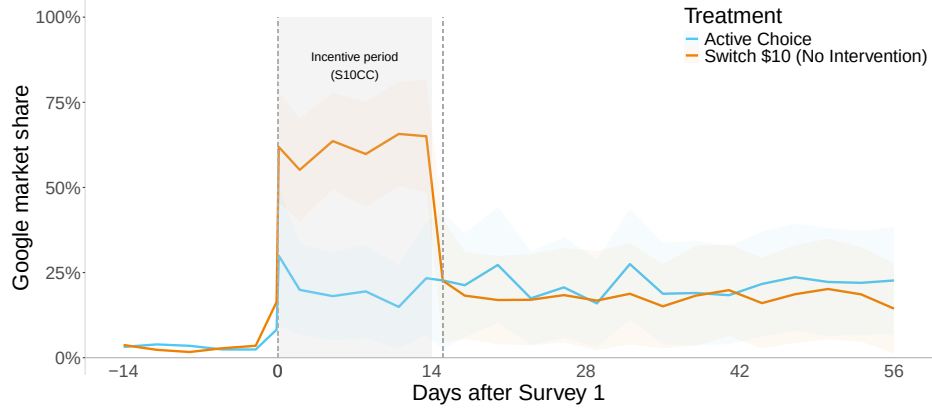
Among baseline Bing users, on the other hand, we find that the Google market share increased from 3.8 percent to 18 percent in the Active Choice group. The larger effect for Bing users is consistent with many permanently inattentive Edge users preferring Google—and thus switching when given an active choice—



(a) Baseline Google Users: Active Choice, Control & Default Change Groups



(b) Baseline Google Users: Switch Bonus Group



(c) Baseline Bing Users

Figure 3: Search Engine Market Shares by Treatment Group

Notes: This figure presents Bing market shares for Baseline Google users in Panels (a) and (b), and Google market shares for Baseline Bing users in Panel (c), broken down by treatment group over time. The dashed vertical lines mark the dates of the two surveys, and the shaded area indicates the incentive period. To compute daily market shares, we first compute the daily market share for each participant and then average shares across participants. To smooth and clarify the figure, we plot averages for groups of two or three days. We define day 15 as the day users take Survey 2, and measure all subsequent days relative to this point, excluding data after day 14 and before Survey 2 completion. In Panel 3a, “Switch (No Intervention)” refers to all users in the S1, S10CC, and S25 groups.

but relatively few permanently inattentive Chrome users preferring Bing. It could also reflect switching costs being larger for Bing users. Either way, our results suggest that mandating choice screens on all devices (including those currently subject to a Bing default) may increase Google’s overall market share.

Switch Bonus Group. We now analyze the market shares of baseline Google users in the Switch Bonus group (Figure 3a). In this section we focus on participants who experienced Bing without Ranking Degradation or Ad Blocking.

During the two-week incentive period, pooling across the different price groups, the Bing market share in this group was 60 percent. Figure 3b breaks out the underlying market shares for the different price levels. We see a clear price response during the incentive period. Even a \$1 payment leads to a 32 percent Bing market share—30 percent higher than for Active Choice—suggesting that an important fraction of users have weak preferences. A \$10 payment increases Bing’s market share to 64 percent, and a \$25 payment to 74 percent. The substantial Google market share, even for higher prices, suggests a fat-tailed distribution of willingness to accept with many people weakly and some people strongly preferring Google. Our survey confirms this: 1.3 percent of users say that Bing is a lot better than Google, while 61 percent of users say that Google is a lot better (Appendix Figure A5 shows the histogram of responses).

Strikingly, many of the users that our payment convinced to try Bing actively chose to continue using Bing after the end of the incentive period. During the week after Survey 2, the Bing market share among those who accepted our offer and qualified for the incentive payment was 38 percent, and it was still 35 percent during the last week of the experiment. This result starkly contrasts with the Active Choice group. The market share of 22 percent for the Switch Group as a whole (including participants who declined our offer) in the last week of the experiment is substantially higher than the 2.5 percent in the Active Choice group. The only difference between these two groups is that Switch Group participants were exposed to Bing for two weeks. Our results hence suggest that their perceptions about Bing improved after exposure.

Ratings and exit survey responses support the hypothesis that participants increased their relative preference for Bing. Table 3 shows the initial ratings in Survey 1 and the relative ratings change from Survey 1 to Survey 2 for participants in the \$10 Switch Group that accepted our \$10 offer, regardless of whether they met the conditions to receive the bonus (Appendix Figure A3 presents average responses in each survey).¹³ A higher relative rating indicates a stronger preference for Bing. On average, Google users’ ratings of Bing significantly improved in all dimensions except the number of ads shown. These changes are sizable. The effect on the difference in the overall quality rating between Google and Bing corresponds to a third of the initial gap and more than half a standard deviation. In our exit survey, when we asked former Google

¹³Table 3 focuses on S10CC users only. Section 4.4 shows how changes in ratings are affected by the Ranking Degradation and Ad Blocking treatments.

Table 3: **\$10 Switch Group: Change in Ratings for Bing Relative to Google**

	(1) Overall quality	(2) Relevant links	(3) Result page features	(4) Relevant ads	(5) AI chat	(6) Privacy	(7) Rewards	(8) Number of ads
Panel A: Baseline Google Users								
Δ Rating	0.403*** (0.074)	0.231*** (0.078)	0.511*** (0.090)	0.452*** (0.081)	0.489*** (0.077)	0.425*** (0.063)	0.353*** (0.066)	-0.077 (0.058)
Rating at baseline								
Mean	-1.475	-1.281	-1.158	-0.851	-0.421	-0.633	0.231	0.683
Sd	0.784	0.849	0.883	0.831	1.078	0.923	1.081	0.719
<i>N</i>	221	221	221	221	221	221	221	221
Panel B: Baseline Bing Users								
Δ Rating	-0.394* (0.213)	-0.121 (0.212)	-0.303 (0.211)	-0.121 (0.161)	-0.303* (0.171)	-0.121 (0.155)	-0.152 (0.152)	-0.091 (0.118)
Rating at baseline								
Mean	0.364	0.121	0.333	0.091	1.030	0.242	1.576	0.848
Sd	1.295	1.244	1.021	0.843	0.883	0.663	1.001	0.712
<i>N</i>	33	33	33	33	33	33	33	33

Notes: This table presents average changes in ratings for Bing relative to Google from Survey 1 to Survey 2. The sample is participants in the \$10 Switch Bonus group that passed the attention check and accepted the offer to switch. We further restrict our sample to users without Ranking Degradation or Ad Blocking. The survey questions in column 1 and columns 2–7, respectively, are “Overall, how would you rate the quality of Google relative to Bing?” and “How would you rate the quality of Google relative to Bing on the following dimensions?” Response options were “Bing is a lot better,” “Bing is a little better,” “They are about the same,” “Google is a little better,” and “Google is a lot better,” coded as 2, 1, 0, -1, and -2, respectively. The survey question in column 8 is “How do you feel about the number of ads on Bing?” Response options were “way too many,” “too many,” “right amount,” “too few,” and “way too few,” coded as 2, 1, 0, -1, and -2, respectively. *, **, ***: statistically significant with 90, 95, and 99 percent confidence, respectively. Standard errors clustered at the participant level.

users to give one or more reasons why they decided to stay with Bing after the incentive ended, 64 percent responded that Bing is better than expected, 59 percent that they got accustomed to it, 5 percent that they always wanted to use it, and 28 percent gave other reasons in a free form text response. The first and second responses above are *both* consistent with our model. Admittedly, they imply different interpretations of why the perceived quality term ζ changes. In the first case, it changes because people learned about Bing, and in the second because the utility of using Bing increased after getting accustomed to it. However, the perceived quality increases in both cases, capturing the fact that users’ perceived utility of using Bing went up. Therefore, our counterfactual results are invariant to these two interpretations as they both affect market shares and consumer surplus in the same way.¹⁴

Among baseline Bing users, the Google market share is 59 percent during the incentive period. We also

¹⁴Another possibility is that during the incentive period, not only does utility for Bing increase, but utility for Google decreases as users get unaccustomed to it. In counterfactuals, this would yield the same market shares as our current model, but smaller effects of Bing exposure on consumer surplus.

find that an important fraction of these users decide to stay with Google: the Google market share was 16 percent after Survey 2. However, unlike what we observe among baseline Google users, this market share is not statistically different from the 18 percent market share among Active Choice users. This suggests that Bing users are well-informed about Google’s quality. The results from the rating survey are mostly consistent with this interpretation (Table 3): Bing users in general update less towards Google, and only the overall quality and AI chat updates are marginally significant.

Default Change. In the Default Change group, Bing’s market share during the two-day incentive period was 73 percent, trending down gradually over the next few weeks: it was 50 percent seven days after Survey 1, and it was 44 percent on day 14.¹⁵ This gradual decline is consistent with our stochastic model of inattention. The market share eventually stabilizes at 40 percent after four weeks, significantly higher than the 17 percent who still use Bing in the Switch Bonus group at that time.

There are two potential explanations for this persistently high Bing market share after the default change. First, some users may be permanently inattentive, as suggested by the fact that the Bing market share in the Default Change group at the end of the experiment is significantly higher than the 17 percent in the Switch Bonus group. Second, users may have changed their perceptions of Bing’s quality during the incentive period, consistent with learning in the Switch Bonus Group. The results from our exit survey suggest that both explanations play a role: while 44 percent of participants reply that they still use Bing because they forgot to switch back or that switching back was too much effort, 35 percent of participants reply that they kept using Bing because they prefer it. The significant share of people who revise their perceptions about Bing after the default change suggests another important effect of Google’s default agreements: it prevents users from learning about the search engine that would otherwise be the default.

4.4 Effects of Ranking Degradation and Ad Blocking

We now measure the effects of Ranking Degradation and Ad Blocking. This analysis is limited to baseline Google users in the \$10 Switch Bonus group—the only group that we randomized into Ad Blocking or Ranking Degradation. We further limit our sample to participants who accepted the Survey 1 offer to switch to Bing for \$10, since users who declined the offer were not exposed to these interventions.

Survey ratings. We first investigate the effects of these interventions on search engine ratings. We regress the change in the ratings from Survey 1 to Survey 2 on indicators for the Ranking Degradation and Ad Blocking groups. We also include a constant, which measures the average change for users who were not subject to Ranking Degradation and Ad Blocking—the effect we measured in Table 3. Panel A of Table

¹⁵We do not observe any meaningful change in the time trend when these users received the placebo Survey 2.

4 presents results. Ranking Degradation significantly reduces the positive updating about Bing on overall quality, the relevance of links, the result page features, and the relevance of ads. Interestingly, Ranking Degradation’s effect on the relevance of results is similar to (though a little larger than) the learning effect reported in Table 3 and implied by the constant. This indicates that Google users in Survey 1 expect the results of Bing to be (almost) as bad as we make them with Ranking Degradation. Ad Blocking significantly reduces the positive update about the relevance of ads and the result page features.

Click-through rates. We next investigate the interventions’ effect on Bing results page clicks. To do this, we regress different measures of click-through rates—a standard industry measure of result-page relevance—on indicators for the Ranking Degradation and Ad Blocking groups.

Panel B of Table 4 presents results. The first two specifications are manipulation checks that show our Ranking Degradation and Ad Blocking treatments worked as expected. Columns 3–5 show the effects of Ranking Degradation on various click-through rates, which reflect whether participants find the results useful and will also serve as our measure of result relevance in the returns-to-scale analysis. Our results suggest that participants find the organic results after Ranking Degradation less useful: the overall click-through rate on organic links drops by 7.51 percentage points (Column 4), and the top link click-through rate falls by 15.3 percentage points (Column 5). Instead, participants are more likely to click on ads, which attenuates the overall effect on clicks (Column 3). Ad Blocking, on the other hand, does not significantly affect participants’ interaction with organic search results, nor does it much affect overall click-through rates.

Market shares. Figure 4 shows the effects of the Ranking Degradation and Ad Blocking treatments on Bing market shares. Let Y_{it} be Bing’s market share, with t indexing periods specifically defined for this regression (and pre-registered as part of our pre-analysis plan), and denote indicators for the Ranking Degradation and Ad Blocking groups by w_i^R and w_i^A respectively. We estimate the following regression:

$$Y_{it} = \tau_t^R w_i^R + \tau_t^A w_i^A + \mu_t + \varepsilon_{it}. \quad (8)$$

Figure 4 shows the τ_t^R and τ_t^A coefficients. Both interventions reduce Bing’s market share after the active choice in Survey 2. These effects, however, are not statistically significant, despite the fact that the extension intervention changes participants’ behavior on the result page as well as their relative rating of Bing versus Google. This result suggests that, when making choices, participants simply do not place enough weight on these attributes relative to other considerations, such as the interface.¹⁶

¹⁶For a limited time, a malfunctioning of our extension may have affected users in the ad-block treatment who clicked beyond the first result page. While removing the potentially affected users (in a way that is balanced across groups) does not alter statistical

Table 4: Effects of Ranking Degradation and Ad Blocking on Quality Ratings and Bing Clicks

Panel A: \$10 Switch Group: Change in Ratings of Bing Relative to Google								
Dep. var:	(1) Overall quality	(2) Relevant links	(3) Result page features	(4) Relevant ads	(5) AI chat	(6) Privacy	(7) Rewards	(8) Number of ads
Ad Blocking	-0.095 (0.089)	-0.060 (0.090)	-0.192* (0.106)	-0.206** (0.090)	-0.108 (0.092)	-0.026 (0.073)	-0.135* (0.079)	0.119 (0.075)
Ranking Degradation	-0.374*** (0.090)	-0.403*** (0.091)	-0.406*** (0.107)	-0.260*** (0.091)	-0.103 (0.092)	-0.128* (0.074)	0.062 (0.079)	0.010 (0.074)
Constant	0.537*** (0.081)	0.264*** (0.083)	0.714*** (0.093)	0.556*** (0.084)	0.630*** (0.081)	0.470*** (0.064)	0.458*** (0.067)	-0.090 (0.063)
R ²	0.028	0.031	0.027	0.021	0.004	0.005	0.005	0.004
N	646	646	646	646	646	646	646	646

Panel B: Clicks On Bing					
Dep. var:	(1) Original search rank	(2) CTR ads	(3) CTR all	(4) CTR organic	(5) CTR organic (top)
Ad Blocking	0.033 (0.096)	-0.026*** (0.003)	-0.013 (0.027)	0.012 (0.026)	-0.017* (0.010)
Ranking Degradation	3.422*** (0.090)	0.019*** (0.003)	-0.057** (0.027)	-0.075*** (0.027)	-0.153*** (0.010)
Constant	1.640*** (0.054)	0.017*** (0.002)	0.361*** (0.017)	0.343*** (0.017)	0.237*** (0.008)
R ²	0.666	0.131	0.008	0.014	0.294
N	646	636	636	636	636

Notes: This table presents estimates of the effects of Ranking Degradation and Ad Blocking on ratings of Bing relative to Google (panel A) and Bing click outcomes (panel B). The sample is baseline Google users in the \$10 Switch Bonus group who accepted the offer to switch to Bing on Survey 1. In Panel A, the survey questions in column 1 and columns 2–7, respectively, are “Overall, how would you rate the quality of Google relative to Bing?” and “How would you rate the quality of Google relative to Bing on the following dimensions?” Response options were “Bing is a lot better,” “Bing is a little better,” “They are about the same,” “Google is a little better,” and “Google is a lot better,” coded as 2, 1, 0, -1, and -2, respectively. The survey question in column 8 is “How do you feel about the number of ads on Bing?” Response options were “way too many,” “too many,” “right amount,” “too few,” and “way too few,” coded as 2, 1, 0, -1, and -2, respectively. In panel B, the outcomes are the average original rank (before ranking degradation) of the Bing organic results the user clicked on (column 1), the Bing click-through rate (CTR) including only ad clicks (column 2), Bing CTR including both ad clicks and organic result clicks (column 3), Bing CTR including only organic result clicks (column 4), and Bing CTR including only the clicks associated with the first-ranked Bing search result (column 5). *, **, ***: statistically significant with 90, 95, and 99 percent confidence, respectively. Standard errors clustered at the participant level.

5 Model Estimation

We now explain how we use our experiment to estimate the model. First, we describe how we map the data to the model. Next, we detail how the experimental treatments identify the model parameters and outline significance, it attenuates the estimated market share effect of Ad Blocking.

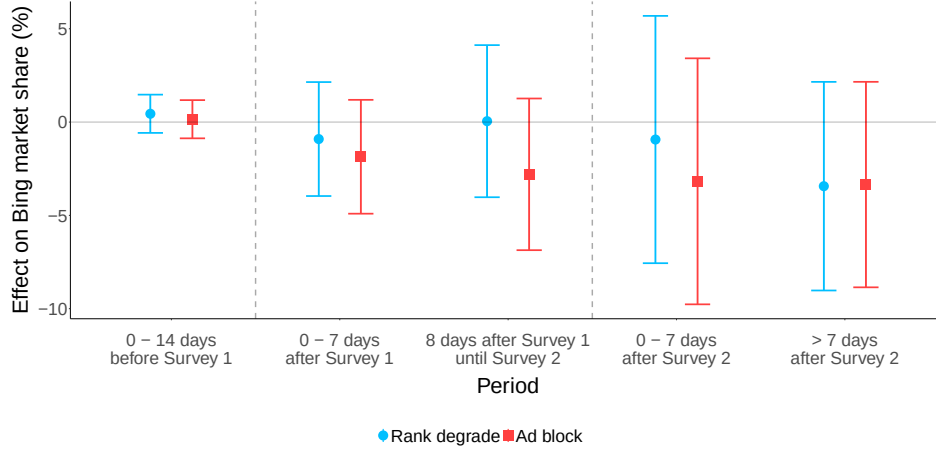


Figure 4: **Effects of Ranking Degradation and Ad Blocking on Market Share**

Notes: This figure presents the effects of Ranking Degradation and Ad Blocking on the Bing market share over time, estimated using equation (8). The sample includes only baseline Google users that accepted the \$10 offer to switch to Bing for the 14 days between Survey 1 and Survey 2. Whiskers indicate 95 percent confidence intervals derived from standard errors clustered at the participant level. Dashed vertical lines mark the dates of the two surveys.

the generalized method of moments procedure we use for estimation.

We let t index two-week periods. Period $t = 0$ is the 14 days before Survey 1, $t = 1$ is the 14 days between Survey 1 and Survey 2, and $t = 2, 3, \dots$ are successive 14-day periods after Survey 2.

In the model, consumers make a binary choice between Bing and Google. In the data, a small number of people multi-home or occasionally use other search engines. For estimation, we say those users chose the search engine where they conducted the most searches.

Our model predicts browser-specific market shares, but our experimental treatment assignments depend on baseline search engine. For instance, baseline Bing users are paid to switch to Google in the Switch Bonus condition, and baseline Google users are paid to switch to Bing. To simplify the computation of market shares, we make a natural assumption that people who had already switched away from their browser’s default search engine at baseline (e.g., Chrome users who search on Bing) would also do so if paid an incentive. This allows us to compute market shares at the browser level, as in our model.

5.1 Identification

We now provide an intuitive discussion of identification using Figure 5, which provides a stylized representation of how market shares evolve over time for Chrome users. The arguments trivially extend to Edge users. We present formal identification arguments in Appendix B.1.

The perceived quality difference $\Delta \tilde{\zeta}$ is identified by the market share in the Active Choice group. Since these users must make a choice without any defaults or switching costs influencing them, a larger Bing share

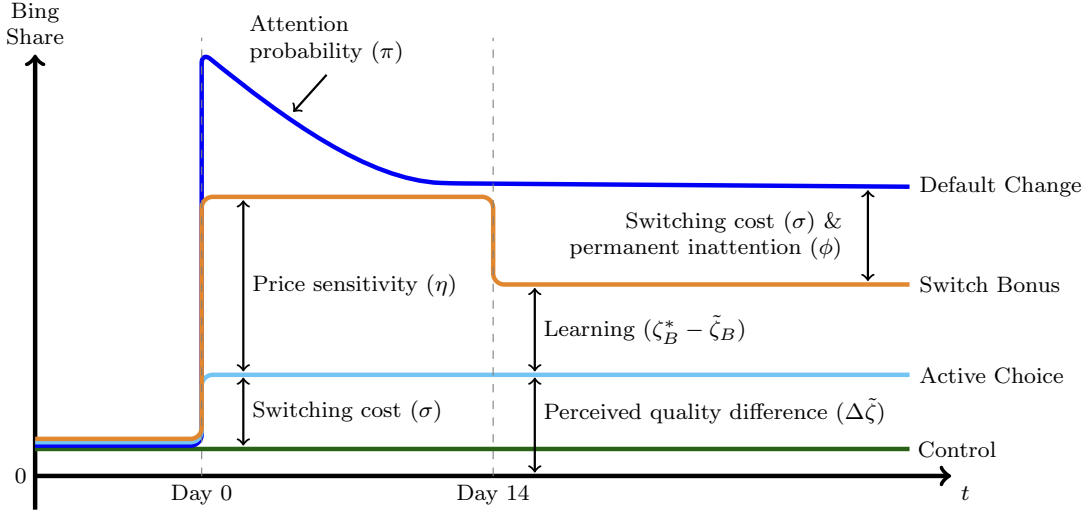


Figure 5: **Relation Between Experimental Market Shares and Parameters**

Notes: This figure illustrates, for Chrome users, the Bing market shares over time in each treatment condition. It also shows which elements identify each of the model parameters.

indicates higher perceived quality of Bing relative to Google.

The price sensitivity η is identified by the difference in market shares between the Switch Bonus group and the Active Choice group during the incentive period. Both groups make an active choice in Survey 1, but only the Switch Bonus group is paid.

Learning $\zeta_B^* - \tilde{\zeta}_B$ is identified by the market share difference between the Switch Bonus group and the Active Choice group after the incentive expires on day 14. When this incentive expires, users in the Switch Bonus group are asked to make another active choice. Hence, any subsequent difference in market shares between these two groups must be driven by the Switch Bonus group's additional experience with Bing.

The attention probability π is identified by the decaying pattern in the Bing market share of the Default Change group: the higher π , the more rapidly this market share will decay towards its long-run level.

The switching cost σ and the share of permanently inattentive users ϕ are jointly identified by two market share differences: the long-run difference between the Default Change group and the Switch Bonus group, and the difference between the Active Choice group and the Control group. Both differences are larger when these two sources of inertia are important. Without inertia, Control Group users would already be with their preferred search engine (as in the Active Choice group). Similarly, without inertia, all Default Change users would eventually switch to their preferred search engine—just like Switch Bonus do after Survey 2.

While formally both σ and ϕ affect both differences, in practice the difference between Active Choice and Control is almost entirely driven by the switching cost σ . Intuitively, permanent inattention ϕ only

affects users wanting to switch. Given Bing’s small market share in the Active Choice group (2.6%), few Control users would like to switch. By contrast, Default Change users only switched to Bing because of the incentive payment, and therefore many of them would likely want to switch back. Hence, permanent inattention affects the difference between Default Change and Switch Bonus but has almost no effect on the difference between Active Choice and Control. Intuitively, this allows us to recover switching cost σ from the difference between Active Choice and Control. Then, once σ is known, what remains to be explained about the gap between Default Change and Switch Bonus informs us about the extent of permanent inattention ϕ .

The distribution of idiosyncratic preferences $\Delta\chi_i$ is identified from market shares across various Switch Bonus group payments and the Active Choice group: as we increase the incentive, we observe what fraction of users switches to Bing, measuring quantiles of the idiosyncratic preference distribution (illustrated in Figure A8, Appendix B.1). With sufficiently many prices, the distribution could be identified non-parametrically; since our experiment includes only four prices, our estimation uses a parametric assumption.

5.2 Estimation

We estimate our model parameters by the generalized method of moments (GMM). The parameters that we estimate are the perceived quality difference $\Delta\tilde{\zeta}$, learning $\zeta_{-d}^* - \tilde{\zeta}_{-d}$, the ad load response α , the relevance response ρ , the price response η , the amortized switching cost $\sigma(1 - \delta)$, the fraction of inattentive consumers ϕ , the per-period probability of paying attention among attentive consumers π , and parameters for the distribution of idiosyncratic preferences.

We assume the distribution of idiosyncratic preferences follows a shifted lognormal distribution, which allows for more flexibility than commonly used distributions, such as normal or logistic. After normalizing its mean and variance, this distribution has one remaining parameter $\gamma \in (0, \infty)$ that captures its skewness. It thus ranges from a symmetric distribution when $\gamma = 0$ to a heavily skewed distribution as $\gamma \rightarrow \infty$, with a fat tail of customers that have a strong preference for Google—as suggested by the fact that many users are unwilling to switch to Bing for a payment of \$25.¹⁷ With our normalization, the preference shifters $\Delta\tilde{\zeta}$ and $\Delta\zeta^*$ represent differences from the value that would result in a market share of one half.

We refer to the full vector of parameters as θ . Theoretically, we can identify each element of θ separately for Edge and Chrome users. In practice, our small sample of Edge users forces us to pool estimation of most parameters to economize on power. Hence, we only allow the quality difference $\Delta\tilde{\zeta}$ and learning $\zeta_{-d}^* - \tilde{\zeta}_{-d}$ to differ between Chrome and Edge users. Since Figure 3c and Table 3 show that Bing users do not learn much about Google, and since the relevant market shares are noisy, we directly set $\zeta_G^* = \tilde{\zeta}_G$. We set the rest

¹⁷We normalize the distribution so that $S(-1) = 0$ and $S(0) = 1/2$. The PDF of the normalized distribution of $\Delta\chi$ is $\exp(-\log(x+1)^2/(2\gamma^2)) / (\sqrt{2\pi}\gamma(1+x))$.

of the parameters to be the same for all users, no matter whether they are on Edge or on Chrome.

The moments that we use closely follow the identification arguments from Section 5.1. All moments are based on market shares that we observe in the data, which we present in Appendix Table A6. The first set of moments are simply market shares: the baseline market share s_{-d0} , the Active Choice market share $s_{-d,t \geq 1}^A$, the market shares for the Switch Bonus group during the incentivized period at different prices $s_{-d,t=1}^{S1}$, $s_{-d,t=1}^{S10CC}$, and $s_{-d,t=1}^{S25}$, and the post-Survey 2 market shares of the Switch Bonus group under different interventions $s_{-d,t \geq 2}^{S10CC}$, $s_{-d,t \geq 2}^{S10RC}$, $s_{-d,t \geq 2}^{S10CA}$, and $s_{-d,t \geq 2}^{S10RA}$. Following the arguments in Section 5.1, these shares identify the distribution of idiosyncratic preferences, perceived differences in quality, learning, the price response, and quality preferences. To identify the attention probability π , we exploit the market shares of the Default Change group right after Survey 1, after one week, and after two weeks. Based on these three shares, we write out a moment condition that captures how quickly the market share converges to its long-run value (see Appendix B.2). Finally, to identify switching costs and inattention, we need moments for $s_{-d,t \geq 0}^A - s_{-d,t \geq 0}^C$ and $s_{-d\infty}^D - s_{-d,t \geq 2}^{S10CC}$. We already included moments corresponding to s_{-d0} (which is the same as $s_{-d,t \geq 0}^C$), $s_{-d,t \geq 1}^A$, and $s_{-d,t \geq 2}^{S10CC}$, so we include an additional moment that characterizes $s_{-d\infty}^D$ as a function of shares that we observe in the data. See further details in Appendix B.2.

Our estimates are given by $\hat{\theta} = \text{argmin}_{\theta} G(\theta)' \Omega G(\theta)$ where m indexes different moments, $G_m(\theta)$ is the average of $g_{im}(\theta)$ across users, and Ω is a weighting matrix. We use two-step GMM.¹⁸

5.3 Demand Parameter Estimates

Table 5 presents the demand parameter estimates. Figure 6 shows that the distribution of idiosyncratic preferences is very skewed, which fits well the price responses described in Section 4.3: many users are close to indifferent between both search engines, and there is a fat tail of users that strongly prefer Google. This has two important implications. First, a given price or quality change has a large market share impact when the marginal users are close to indifferent (when Google’s market share is high), but only a modest impact when the marginal users belong to the fat tail of users with strong preferences in favor of Google (when Google’s market share is low). This feature drives many of our findings below. Second, the combination of a fat tail and permanent inattention implies large consumer surplus losses due to agents with strong preferences that nevertheless use an undesired search engine because they are inattentive. Hence, when computing consumer surplus, we censor idiosyncratic preferences at \$25 to avoid our results being driven by extrapolation beyond the price offers that participants received.

We now discuss parameters for Chrome users as well as the rest of the parameters for all users—which are mainly identified from data on Chrome users because they account for most participants in our sample.

¹⁸We first set $\Omega = \text{Cov}(G(\theta))^{-1}$ for arbitrary θ , and then we set Ω to be the optimal weighting matrix given the initial estimates.

Table 5: **Demand Parameter Estimates**

Description	(1) Formula	(2) Estimate	(3) SE	(4) 95% CI	
All users					
Distribution shape	γ	2.90	0.68	[2.15	4.33]
Price response	η	0.31	0.13	[0.22	0.55]
Permanent inattention	ϕ	0.34	0.07	[0.20	0.41]
Attention probability	π	0.82	0.19	[0.38	0.97]
Amortized switching cost	$\sigma(1 - \delta)/\eta$	\$0.004	0.060	[-\$0.005	\$0.034]
Baseline Chrome users					
Bing preference shifter	$\Delta\tilde{\zeta}/\eta$	-\$3.20	0.81	[-\$4.49	-\$1.85]
Learning	$(\zeta_{-d}^* - \tilde{\zeta}_{-d})/\eta$	\$0.21	0.17	[\$0.04	\$0.58]
Ad load response	$\alpha(a_{-d}^{CA} - a_{-d}^{CC})/\eta$	-\$0.04	0.07	[-\$0.16	\$0.05]
Quality response	$\rho(r_{-d}^{RC} - r_{-d}^{CC})/\eta$	-\$0.02	0.07	[-\$0.13	\$0.07]
Baseline Edge users					
Google preference shifter	$\Delta\tilde{\zeta}/\eta$	-\$7.56	0.68	[-\$8.84	-\$6.57]
Learning	$(\zeta_{-d}^* - \tilde{\zeta}_{-d})/\eta$	\$0	-	-	

Notes: This table presents the parameter estimates from the demand estimation procedure described in Section 5. All utility parameters have been expressed in units of dollars per two-week period by dividing through by the price response η . Standard errors and confidence intervals are bootstrapped by resampling participants.

We use our estimate of the price response η to interpret all other parameters in units of dollars per two-week period. For Chrome users who have not used Bing, we observe a negative Bing preference shifter $\Delta\tilde{\zeta}$ (equivalent to a payment of \$3.20 per two-week period), consistent with the low Bing market share in the data. After exposure to Bing, its perceived utility increases by \$0.21 per two-week period. Updating perceptions about Bing affects those users who are close to indifferent between both search engines, where there is a large density of users. Thus, as we note above, a small change in the perception of quality is enough to generate a large change in market shares.

The estimated switching costs $\sigma(1 - \delta)$ are small (0.4 cents), given the small effect on market shares of the Active Choice intervention—which, as we explain in Section 5.1, are mainly driven by switching costs rather than inattention. On the other hand, we find that inattention plays an important role as 34 percent of users are permanently inattentive. The per-period attention probability π of 82 percent means that users who are not permanently inattentive make attentive choices frequently. Our estimates for the preferences for quality indicate that Ad Blocking decreased utility as much as a price decrease of \$0.04, and Ranking Degradation caused a utility decrease of \$0.02. Neither of these two parameters is statistically significant.

For Edge users, we find a negative Google preference shifter $\Delta\tilde{\zeta}$ that is equivalent to a payment of \$7.56 per two-week period, consistent with the high share of Bing among Edge users.

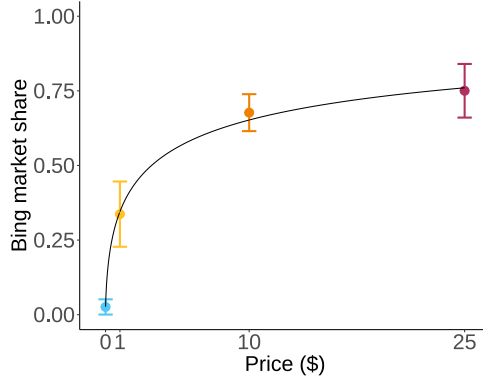


Figure 6: **Distribution of Payment Required to Switch to Bing**

Notes: This figure shows our estimates of the fraction of users willing to switch their default search engine to Bing for two weeks (on the vertical axis) for a given payment (on the horizontal axis). A price of zero represents the Active Choice group, and the other prices represent different payments within the Switch Bonus group. The \$10 price only includes users in the S10CC group. We exhibit means with associated standard errors (dots and associated error bars) and the distribution that we fit to these market share moments (the solid line). The plot confirms that our parametric assumption fits the data well.

6 Economies of Scale in Data

In this section, we estimate the extent to which collecting additional data enables search engines to present more relevant results. To that end, we use data on search terms and clicks from Bing. As is standard industry practice, we use the probability that a user clicks on the top-ranked result as our measure of the relevance of the links presented on a search results page. Henceforth, we will refer to search terms as queries and instances where a user enters a query as searches.

Google argues that the relationship between data and search result relevance exhibits quickly diminishing returns to scale, implying that increased scale has little effect on the relevance of results (e.g., [Varian 2015](#)). This argument is less convincing if there are many *tail* queries (i.e., queries with few searches) for which additional data may still be valuable. Indeed, examining all searches made on Bing over 12 months in 2021 and 2022, we find that the distribution of queries exhibits a long tail: more than 38.7 percent of searches are for rare queries that are searched less than 100 times.¹⁹ For this reason and for reasons of identification that we explain below, we will focus on new queries, which have not received too many searches.

For our analysis, we randomly sampled 43,991 new search terms, i.e., search terms for which Bing had no search record between January 2021 and January 2022. For these search terms, we recorded each search in the subsequent year (i.e., between January 2022 and January 2023).²⁰

¹⁹Appendix Figure A9 shows what fraction of searches accrues to queries with different occurrence rates.

²⁰The resulting panel dataset at the (query, search)-level records (i) the date of each search, (ii) various click-related outcome measures (was there any click, was the top result clicked, was there any click from which the user did not immediately return) for each search, and (iii) an identifier for the URL that was top-ranked. We report summary statistics in Table A7 in Appendix C.1. Optimally, we would subset to queries Bing has never seen, but we are limited by Bing’s retention period of 24 months.

6.1 Empirical Strategy

Our goal is to measure the causal effect of the number of searches n_{qt} on result relevance r_{qt} for query q at time t :

$$r_{qt} = f(n_{qt}; \theta) + \varepsilon_{qt}.$$

We need to address two potential forms of confounders. First, there could be persistent differences in the difficulty of serving results across queries. For instance, navigational queries (“facebook.com”) are both common and easy to answer. Thus, a simple cross-sectional comparison may bias results in favor of finding returns to scale. To control for such cross-query confounders, we use previously unseen queries over time as they gather their first searches and control for query fixed effects (as in [He et al., 2017](#)).

However, this strategy does not address time-varying confounders, which may arise due to compositional changes in the types of users over the lifetime of a query. For instance, suppose a query is concerned with a news event. For users who arrive early in the query’s lifetime, the event may be more newsworthy, and they are hence more likely to click on results. By contrast, users that arrive later may already know of the event and just want to verify its date without the need to click. To deal with such time-varying confounders, we use an identification strategy that only uses variation in click-through rates that arises from the order in which search engines present results. [Figure 7](#) presents event studies that highlight this variation. It displays click-through rates against the number of searches since the first or second time its top-ranked result changed. Updates to the ranking increase click-through rates, suggesting the search engine learns to serve more useful results over time. We also observe a general downward trend unrelated to search results.

To obtain a causal estimate from this variation, we rely on a key identifying assumption, which seems plausible in this context: any *causal* effect of more search data on clicks has to be due to the results that the search engine serves and the order in which it presents them.²¹ In [Appendix C.3](#), we show formally that, under this assumption, we can exploit the front door criterion ([Pearl and Mackenzie 2018](#); [Imbens 2020](#); [Bellemare, Bloem, and Wexler 2024](#)) to purge confounding variation in click-through rates.

6.2 Estimation

Building on the aforementioned identification argument, we now estimate the relationship between search result relevance (as proxied by the click-through rate) and the number of prior searches for a given query q on date t . For simplicity, we specify a functional form that allows for different relationships between data

²¹This assumption would be violated, for example, if the search engine also uses data to improve its interface. However, to cause concern for our identification strategy, such interface adjustments would have to be query specific, which seems less plausible.

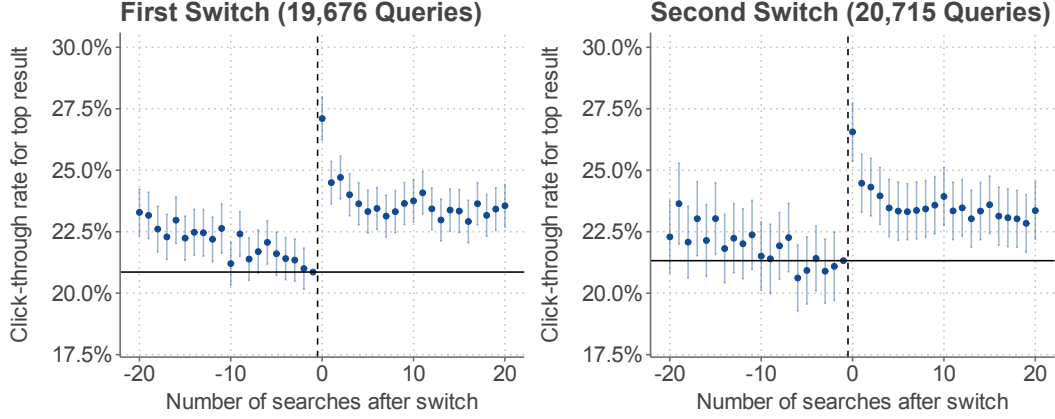


Figure 7: **Event Study Illustrating Effect of Ranking Change on Click-Through Rates**

Notes: We regress the CTR on query fixed effects and a set of dummies measuring how many searches happened since the k -th switch in the top-ranked result. These figures report the coefficient on these dummies, providing evidence that CTR increases at the time of a ranking change but otherwise follows a secular decline. The secular decline motivates our identification strategy (as it suggests the presence of a potential confounder), and the positive effect of ranking changes suggests that rankings do affect CTR. This Figure restricts attention to isolated switches, i.e., switches for which there is no additional switch within twenty searches of the original switch.

and CTR, including linear ($\theta = 0$), logarithmic ($\theta = 1$), and more concave than logarithmic ($\theta > 1$):

$$r_{qt} = \alpha_q + \beta \frac{(n_{qt})^{1-\theta} - 1}{1-\theta} + \varepsilon_{qt}. \quad (9)$$

As we show below, it fits the data quite well.

To ensure that we only use CTR variation deriving from ranking changes, we estimate equation 9, but we replace CTR r_{qt} with $\hat{r}_{qt}$, a prediction of the CTR that is based only on the top result. Appendix C.3 clarifies that this estimator amounts to measuring a causal effect using the front-door criterion (Pearl and Mackenzie 2018). See Appendix C.4 for estimation details. Since our equilibrium model does not distinguish between different queries, we choose an overall intercept α by matching the average click-through rate predicted by our model to that in our experimental data.²²

Table 6 presents the resulting estimates. We can strongly reject the null hypothesis that data does not matter ($\beta = 0$). As to the shape of the returns from data, the estimates strongly point towards a log-linear relationship between CTR and number of searches. Thus, doubling the amount of searches will lead to a fixed increase in CTR, no matter the starting point. Our estimates imply that if Bing had access to Google’s data, CTR would increase from 23.5 percent to 24.8 percent. We show the model’s fit in Figure 8, comparing model predictions to a binscatter plot of the data. The model fits well across many orders of magnitudes.²³

²²This average is computed by integrating over the complete query frequency distribution and weighting each query by its number of searches. While we would optimally use lifetime query frequencies (i.e., query q has seen n views since inception), we actually observe query frequencies over a one year period.

²³He et al. (2017, Fig. 4 and 5) find that going from 300 to 1,900 searches yields a CTR improvement of about 2 percentage

Table 6: **Economies of Scale Estimates**

Description	(1) Parameter	(2) Estimate	(3) SE
Click-through rate at inception	α	0.1811	-
Value of additional data	β	0.0056	(0.0007)
Shape of returns from data	θ	0.9458	(0.0318)

Notes: This table presents our parameter estimates for the relationship between the amount of data and the relevance of search results, as measured by the click-through rate. See equation (9) for the functional form specification.

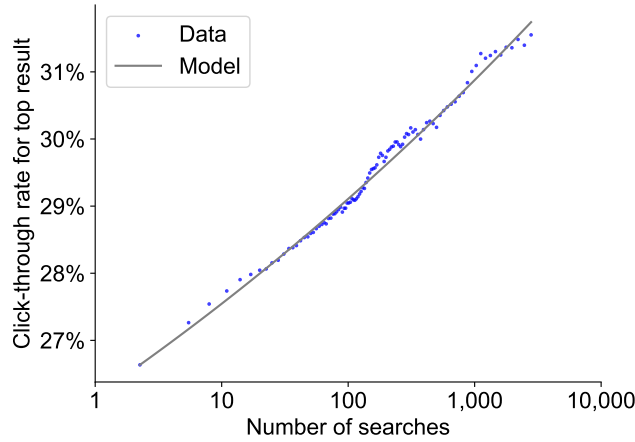


Figure 8: **Model Fit for Economies of Scale**

Notes: We present the model fit for equation (9) by exhibiting both the model-predicted CTR as a function of the number of prior searches for a query as well as a binscatter of the actually observed CTR as a function of the number of prior searches. The model estimates a linear relationship between CTR and the logarithm of the number of prior searches, and this relationship fits the observed data well.

We can use our fitted model to calculate counterfactual average click-through rates on Bing if its market share were to be multiplied by λ : the counterfactual click-through rate $\tilde{r}(\lambda)$ after such an increase in market share is given as a function of the status-quo click-through rate $\bar{r}$ by

$$\tilde{r}(\lambda) = \alpha - \frac{\beta}{1 - \theta} + \lambda^{1-\theta} \cdot \left(\bar{r} - \alpha + \frac{\beta}{1 - \theta} \right). \quad (10)$$

We use this expression in the computation of equilibrium counterfactuals below.

Our approach to estimating economies of scale has several limitations. First, our model is estimated exclusively off new queries. This could be problematic if the effect of data differs between new queries and other types of queries. Second, extending our model to queries with many views requires extrapolation as new queries typically do not reach such high levels of popularity. Still, the fit of our model mitigates this concern. Lastly, the results above do not allow for cross-query learning. In Appendix C.6 we incorporate cross-query learning effects with quantitatively similar results.

7 Counterfactuals

To compute counterfactuals, we adjust the fraction of Chrome and Edge users from our experimental sample to reflect the observed US desktop market shares, decreasing Chrome’s share and increasing Edge’s share.

Direct Effects We first analyze the direct impact of different scenarios on demand without accounting for returns to scale from data. These results represent short-run effects before result relevance changes in response to shifting market shares.

Panel A of Table 7 shows our results. Each counterfactual is compared to a *Status Quo* scenario, in which demand behaves as in the control group, and market shares are given by $s_{-d} = (1 - \phi)S(\Delta\tilde{\zeta} - \sigma(1 - \delta))$. For each scenario, we report consumer surplus and market shares. We understand market shares as an important proxy for overall welfare as they affect both competition for advertisers and quality investment decisions. While these effects are unmodeled, their likely presence means that policies balancing market shares without substantially reducing consumer surplus are probably desirable. Columns 1 and 2 present aggregate market results, while Columns 3-6 break these down by Chrome and Edge users.

To decompose the economic effect of different frictions, we first examine benchmark counterfactuals that are unattainable through realistic policies. Our first scenario, *No Frictions*, eliminates all demand-side frictions: users are fully attentive, perfectly informed about search engine quality, and face no switching costs. Market shares thus solely depend on true quality, $s_{-d} = S(\Delta\zeta^*)$. In this case, Bing’s market share rises

points. Our estimates imply an improvement of 1.5 percentage points. [Schaefer and Sapi \(2023, Fig 5\(c\)\)](#) find that going from 1 to 5,000 searches yields a CTR improvement of 3-5 percentage points. Our estimates imply an improvement of 6.1 percentage points.

Table 7: **Counterfactuals**

Description	Combined		Chrome		Edge	
	(1)	(2)	(3)	(4)	(5)	(6)
	Google share (%)	CS gain (\$/user-year)	Google share (%)	CS gain (\$/user-year)	Google share (%)	CS gain (\$/user-year)
Panel A: Direct Effects						
<i>Benchmarks</i>						
Status Quo	88.9	0.00	98.7	0.00	22.4	0.00
No Frictions	75.8	5.92	82.0	0.61	33.8	41.73
Active Choice	89.1	5.47	97.3	0.08	33.8	41.73
Correct Perceptions	79.7	0.32	88.2	0.36	22.4	0.00
<i>Policy Interventions</i>						
Choice Screen	87.6	0.07	97.3	0.08	22.4	0.00
Bing Default	50.1	-72.93	54.2	-83.75	22.4	0.00
Bing Default + Delayed Choice Screen	74.1	0.06	81.8	0.07	22.4	0.00
Bing Payments (\$10)	51.7	107.36	56.8	92.22	17.2	209.36
Panel B: Equilibrium Effects						
<i>Benchmarks</i>						
Status Quo	88.9	0.00	98.7	0.00	22.4	0.00
No Frictions	75.7	5.92	81.9	0.61	33.8	41.75
Active Choice	89.1	5.47	97.3	0.08	33.8	41.73
Correct Perceptions	79.7	0.32	88.2	0.36	22.4	0.02
No Frictions + Data Sharing	75.6	5.94	81.8	0.62	33.8	41.79
<i>Policy Interventions</i>						
Choice Screen	87.6	0.07	97.3	0.08	22.4	0.00
Bing Default	50.0	-72.91	54.1	-83.74	22.4	0.04
Bing Default + Delayed Choice Screen	74.2	0.06	81.9	0.06	22.4	0.02
Bing Payments (\$10)	51.7	107.37	56.8	92.23	17.2	209.40
Data Sharing	88.9	0.01	98.7	0.00	22.4	0.07

Notes: This table presents counterfactual simulation results. Panel A presents direct effects, before taking into account changes in quality due to economies of scale in data. Panel B presents equilibrium effects, which endogenize quality by accounting for economies of scale in data. Within each panel, the top rows present hypothetical counterfactuals that serve as benchmarks. The bottom rows present counterfactuals that represent policy interventions.

notably from 11 percent to 24 percent, though Google’s share remains large. Consumer surplus increases by \$5.92 per user-year. Although infeasible in practice, this frictionless scenario provides a benchmark for the socially optimal outcome in the absence of scale economies.

Our next two counterfactuals separate the effects of quality misperceptions from inertia. In the *Active Choice* scenario, we shut down switching costs and inattention, so demand matches the Active Choice Treatment ($s_{-d} = S(\Delta\tilde{\zeta})$). Surprisingly, Google’s market share slightly increases by 0.2 percentage points: some Chrome users switch to Bing, but slightly more Edge users switch to Google. Despite minimal changes in shares, *Active Choice* raises consumer surplus significantly (92 percent of the gain under *No Frictions*), as more users choose their preferred search engine. This shows that inertia strongly affects consumer surplus, but does not alone explain Google’s dominant market share. In *Correct Perceptions*, users know the true quality of both search engines but still face switching costs and inattention, with shares given by $s_{-d} = (1 - \phi)S(\Delta\zeta^* - \sigma(1 - \delta))$. Google’s market share drops by 9.2 percentage points, which emphasizes the importance of correcting consumers’ misperceptions. Importantly, the reduction in Google’s market share under *No Frictions* is greater than the sum of the declines from *Active Choice* and *Correct Perceptions*. This is so because most users’ choices are overdetermined in the Status Quo: Google is their default, yet they would choose it anyway based on its perceived quality.

Next, we evaluate proposed policy interventions. The *Choice Screen* scenario, mirroring the European Commission’s 2018 decision on Android choice screens, makes Chrome users actively choose their search engine (as in *Active Choice*) while Edge users remain as in *Status Quo*. Google’s market share declines only slightly (1.3 percentage points), and consumer surplus rises modestly by \$0.07. These results align with [Decarolis, Li, and Paternolli \(Forthcoming\)](#), who found limited market-share effects from choice screens in the EU.

Policies are more successful in moving market shares when they expose many users to the alternative search engine, thus reducing misperceptions. In *Bing Default*, Bing becomes the default search engine across browsers. This approximates proposed bans on Google’s payments for default positions, which would likely result in Bing—the second largest search engine—outbidding other search engines for defaults. Edge demand remains unchanged while Chrome demand shifts towards Bing due to switching costs and inattention, with shares equal to $s_{-d} = \phi + (1 - \phi)S(\Delta\zeta^* + \sigma(1 - \delta))$. Google’s share falls significantly (by 39 percentage points) but consumer surplus declines sharply (\$72.93 per user-year) as permanently inattentive users, including users with strong preferences for Google, are defaulted into a less-preferred option.²⁴

Our counterfactuals suggest that two of the most commonly proposed policies to curb Google’s dominance have important drawbacks. Choice screens increase consumer surplus, but they barely move the

²⁴As explained in Section 5.3, we censor idiosyncratic preferences at \$25, so this value is not driven by the fat tail of preferences. These numbers should thus be thought of as a conservative lower bound for the actual welfare effects.

needle in terms of market shares. Default changes significantly impact market shares, but they come at the cost of a large reduction in consumer surplus. We now consider a policy, *Bing Default + Delayed Choice Screen*, that aims to combine the benefits of both: Bing is initially the default for two weeks, then users actively select their preferred search engine. The initial default allows users to learn about Bing, but the subsequent active choice avoids welfare losses for permanently inattentive consumers who prefer Google.²⁵ This reduces Google’s market share by 15 percentage points with negligible impact on consumer surplus (an increase of \$0.06 per user-year), suggesting that policies can influence market shares without significantly harming consumer surplus. Though challenging to implement, this scenario illustrates the key elements necessary for successful policy—exposing users to other search engines while preserving choice—and demonstrates the magnitude of potential gains.

Finally, we consider a *Bing Payments* counterfactual, in which Bing pays users an additional \$10 (beyond existing rewards) to encourage usage.²⁶ This allows us to determine the extent to which payments could level the playing field in favor of Bing. Shares become $s_{-d} = (1 - \phi)S(\Delta\tilde{\zeta} - \sigma(1 - \delta) + \eta p_B)$, where the payment is $p_B = \$10$ for Chrome but negative for Edge users. This significantly increases Bing’s market share and consumer surplus, though the surplus increase mainly reflects the direct payment to users.

Equilibrium effects We now consider equilibrium effects, incorporating economies of scale in data by endogenizing search result relevance, as detailed in Section 6. Since our estimation of returns to scale in data relies on strong assumptions, these counterfactuals are best viewed as a more speculative exercise.

We first define equilibrium formally. Demand in counterfactual $\mathcal{C}$ is given by:

$$s = D(p, \zeta^*, \tilde{\zeta}; \mathcal{C}), \quad (11)$$

where s denotes market shares, p prices, ζ^* true qualities, and $\tilde{\zeta}$ users’ perceived qualities before learning through experience. This demand function differs across counterfactuals $\mathcal{C}$. For the *Status Quo*, for instance, market shares are given by $s_{-d} = (1 - \phi)S(\Delta\tilde{\zeta} - \sigma(1 - \delta))$, and they are $s_{-d} = S(\Delta\tilde{\zeta})$ for the *Active Choice* counterfactual. See Appendix D for full details.

Based on our model of economies of scale from Section 6, the true quality incorporating scale effects is:

$$\zeta^* = Z(s; \mathcal{C}). \quad (12)$$

²⁵During the first two weeks, market shares match *Bing Default*. Afterward, they shift to mimic *No Frictions* for Chrome users and *Choice Screen* for Edge users. We weight consumer surplus by 2/312 before and 310/312 after the choice screen, approximating a scenario where browsers reset to Bing defaults every six years upon computer replacement. While the sign of the consumer surplus impact depends on the assumed browser reinstallation frequency, its small magnitude is robust.

²⁶Such payments may be hard to implement in practice as people may be tempted to create multiple accounts and perform unnecessary searches to obtain larger payments (Tirole and Biscaglia, 2024).

Without data sharing, $Z(s; \mathcal{C}) = \alpha a_j + \rho \times \tilde{r}(s_j/\hat{s}_j, \bar{r}) + \xi_j$, where $\tilde{r}(\cdot)$ is defined by equation 10 and $\bar{r}$ is the average click-through rate for Bing. This follows from substituting our click-through rate results from Section 6 into quality $\zeta_j^* := \alpha a_j + \rho r_j + \xi_j$. With data sharing, each search engine’s quality is as if they had a market share of one, so $Z(s; \mathcal{C})$ is given by $\zeta_j^* = \alpha a_j + \rho \times \tilde{r}(1/\hat{s}_j) + \xi_j$.

An equilibrium consists of a joint solution to equations 11 and 12 in s and ζ^* . Let $s^{\text{eq}}(\mathcal{C})$ and $\zeta^{\text{eq}}(\mathcal{C})$ denote the equilibrium shares and qualities given counterfactual $\mathcal{C}$.²⁷ From these equilibrium quantities, we derive consumer surplus in equilibrium using the demand function $D(p, \zeta^{\text{eq}}(\mathcal{C}), \tilde{\zeta}; \mathcal{C})$; Appendix D presents all the expressions that we use.

Panel B of Table 7 presents equilibrium counterfactual results. The first rows revisit scenarios from Table 7, now incorporating equilibrium effects. Economies of scale reinforce direct market share effects of interventions by only roughly 1 percent, and their effects on consumer surplus remain similarly limited. This muted effect reflects our findings of modest economies of scale (Section 6) and a weak demand response to quality improvements (Section 4.4).

We also analyze policy scenarios in which regulators require Google to share user data with competitors. This allows Bing to exploit the data from all users, so the Bing result relevance is what it would be if Bing was a monopolist and could hence observe data from all users. In *Data Sharing*, mandated sharing alone has minimal impact: Google’s market share is unchanged after rounding, and consumer surplus rises by just \$0.01. This effect is so small because most users remain unaware of Bing’s quality improvement and thus continue using Google. Even when combined with corrected consumer perceptions (*No Frictions + Data Sharing*), data sharing reduces Google’s market share by less than 0.5 percentage points. The effect is limited for the same reasons why equilibrium effects are small: economies of scale are small, and consumers show limited response to quality.

One caveat for Table 7 is that our estimate of the demand response to the relevance of search results has a wide confidence interval. Hence, Table A9 presents alternative equilibrium results that assume the largest demand response consistent with our estimates. Concretely, we substitute the lower bound of the 95 percent confidence interval for α on Table 5 for its point estimate. Although equilibrium and data-sharing effects become larger, they remain minor relative to scenarios that eliminate demand-side frictions.

8 Conclusion

Google’s large market share in web search is of ongoing concern to both antitrust authorities and regulators. This paper sheds light on this debate, focusing particularly on the role of browser defaults and economies of

²⁷In principle, there could be multiple equilibria: the economies of scale could be so large that the market “tips” towards either Google or Bing. In practice, we measure small economies and a limited demand response to search result relevance.

scale. Our results highlight that browser defaults are partially responsible for Google’s large market share in web search. However, this effect does not arise only because of switching costs and users’ inattention. Our findings show that consumers’ lack of exposure to Bing—partly driven by browser defaults—is a key channel through which Google maintains a higher share than it would have absent any frictions. We also find that sharing Google’s click-and-query data with Microsoft may only have a minor effect on market shares.

Our findings suggest that to significantly shift market shares, regulators must recognize search engines as experience goods and ensure that remedies expose consumers to alternatives. More broadly, our results provide a stark example of how overly pessimistic consumer beliefs about rivals can protect incumbent firms, rendering simple remedies ineffective.

References

- Akerberg, Daniel A. 2003. “Advertising, Learning, and Consumer Choice in Experience Good Markets: An Empirical Examination.” *International Economic Review* 44 (3):1007–1040.
- Agte, Patrick, Claudia Allende, Adam Kapor, Christopher Neilson, and Fernando Ochoa. 2024. “Search and Biased Beliefs in Education Markets.” NBER Working Paper No. 32670.
- Allcott, Hunt, Luca Braghieri, Sarah Eichmeyer, and Matthew Gentzkow. 2020. “The Welfare Effects of Social Media.” *American Economic Review* 110 (3):629–676.
- Allcott, Hunt, Matthew Gentzkow, and Lena Song. 2022. “Digital Addiction.” *American Economic Review* 112 (7):2424–2463.
- Andersen, Steffen, John Y Campbell, Kasper Meisner Nielsen, and Tarun Ramadorai. 2020. “Sources of Inaction in Household Finance: Evidence from the Danish Mortgage Market.” *American Economic Review* 110 (10):3184–3230.
- Aridor, Guy. Forthcoming. “Market Definition in the Attention Economy: An Experimental Approach.” *RAND Journal of Economics*.
- Athey, Susan and Glenn Ellison. 2011. “Position Auctions with Consumer Search.” *Quarterly Journal of Economics* 126 (3):1213–1270.
- Azevedo, Eduardo M, Alex Deng, José Luis Montiel Olea, Justin Rao, and E Glen Weyl. 2020. “A/B Testing with Fat Tails.” *Journal of Political Economy* 128 (12):4614–4672.
- Bajari, Patrick, Victor Chernozhukov, Ali Hortaçsu, and Junichi Suzuki. 2019. “The Impact of Big Data on Firm Performance: An Empirical Investigation.” *AEA Papers and Proceedings* 109:33–37.
- Bellemare, Marc F., Jeffrey R. Bloem, and Noah Wexler. 2024. “The Power of How: Estimating Treatment Effects Using the Front-Door Criterion.” *Oxford Bulletin of Economics and Statistics* 86 (4):951–993.
- Bhargava, Hemant K., Jan Kraemer, and Chayanin Wipusanawan. 2025. “Are Preset Defaults Harmful?”
- Blake, Thomas, Chris Nosko, and Steven Tadelis. 2015. “Consumer Heterogeneity and Paid Search Effectiveness: A Large-scale Field Experiment.” *Econometrica* 83 (1):155–174.
- Blake, Tom, Sarah Moshary, Kane Sweeney, and Steve Tadelis. 2021. “Price Salience and Product Choice.” *Marketing Science* 40 (4):619–636.
- Brot-Goldberg, Zarek, Timothy Layton, Boris Vabson, and Adelina Yanyue Wang. 2023. “The Behavioral Foundations of Default Effects: Theory and Evidence from Medicare Part D.” *American Economic Review* 113 (10):2718–2758.
- Brynjolfsson, Erik, Avinash Collis, and Felix Eggers. 2019. “Using Massive Online Choice Experiments to Measure Changes in Well-being.” *Proceedings of the National Academy of Sciences* 116 (15):7250–7255.

- Brynjolfsson, Erik, Avinash Collis, Asad Liaqat, Daley Kutzman, Haritz Garro, Daniel Deisenroth, and Nils Wernerfelt. Forthcoming. “The Consumer Welfare Effects of Online Ads: Evidence from a 9-Year Experiment.” *American Economic Review: Insights* .
- Bursztyn, Leonardo, Benjamin R Handel, Rafael Jimenez, and Christopher Roth. 2023. “When Product Markets Become Collective Traps: The Case of Social Media.” NBER Working Paper No. 31771.
- Carroll, Gabriel D, James J Choi, David Laibson, Brigitte C Madrian, and Andrew Metrick. 2009. “Optimal Defaults and Active Decisions.” *Quarterly Journal of Economics* 124 (4):1639–1674.
- Chiou, Lesley and Catherine Tucker. 2021. “Search Engines and Data Retention: Implications for Privacy and Antitrust.” In *The Evolution of Antitrust in the Digital Era: Essays on Competition Policy, Volume II*, edited by David S. Evans, Catherine E. Tucker, and Allan Fels. Competition Policy International. Chapter in edited volume.
- Crawford, Gregory S and Matthew Shum. 2005. “Uncertainty and Learning in Pharmaceutical Demand.” *Econometrica* 73 (4):1137–1173.
- DC District Court. 2024. “United States of America et al. v. Google LLC.” Case No. 20-cv-3010 (APM). Accessed at <https://www.texasattorneygeneral.gov/sites/default/files/images/press/Google%20Search%20Engine%20Monopoly%20Ruling.pdf> in March 2025. Memorandum Opinion filed August 5, 2024.
- Decarolis, Francesco, Muxin Li, and Filippo Paternolli. Forthcoming. “Competition and Defaults in Online Search.” *American Economic Journal: Microeconomics* .
- DellaVigna, Stefano and Ulrike Malmendier. 2006. “Paying Not to Go to the Gym.” *American Economic Review* 96 (3):694–719.
- Department of Justice. 2020. “United States vs Google LLC.” Accessed at <https://www.justice.gov/atr/case/us-and-plaintiff-states-v-google-llc> in March 2025.
- Dickstein, Michael J. 2021. “Efficient Provision of Experience Goods: Evidence from Antidepressant Choice.” Working paper.
- Dinielli, David, Fiona M Scott Morton, Katja Seim, Michael Sinkinson, Amelia Fletcher, Gregory S Crawford, Jacques Crémer, Paul Heidhues, Michael Luca, Tobias Salz et al. 2023. “Consumer Protection for Online Markets and Large Digital Platforms.” *Yale Journal on Regulation* 40 (3):875–1125.
- Ecommercedb. 2024. “Customer Journey in Online Shopping: Most Start on Search Engines.” Webpage. Accessed at <https://ecommercedb.com/insights/46-of-retail-customer-journeys-start-with-a-google-search/4465> in Dec 2024.
- Edelman, Benjamin, Michael Ostrovsky, and Michael Schwarz. 2007. “Internet Advertising and the Generalized Second-price Auction: Selling Billions of Dollars Worth of Keywords.” *American Economic Review* 97 (1):242–259.
- Einav, Liran, Benjamin Klopck, and Neale Mahoney. Forthcoming. “Selling Subscriptions.” *American Economic Review* .
- Erdem, Tülin and Michael P Keane. 1996. “Decision-making Under Uncertainty: Capturing Dynamic Brand Choice Processes in Turbulent Consumer Goods Markets.” *Marketing Science* 15 (1):1–20.
- Ericson, Keith M Marzilli. 2014. “Consumer Inertia and Firm Pricing in the Medicare Part D Prescription Drug Insurance Exchange.” *American Economic Journal: Economic Policy* 6 (1):38–64.
- European Commission. 2018. “CASE AT.40099, Google Android.” Accessed at https://ec.europa.eu/competition/antitrust/cases/dec_docs/40099/40099_9993_3.pdf in June 2024.
- Farrell, Joseph and Paul Klemperer. 2007. “Coordination and Lock-In: Competition with Switching Costs and Network Effects.” In *Handbook of Industrial Organization*, vol. 3, edited by Mark Armstrong and Robert Porter. Elsevier, 1967–2072.
- Farronato, Chiara, Andrey Fradkin, and Chris Karr. 2024. “Webmunk: A New Tool for Studying Online Behavior and Digital Platforms.” NBER Working Paper No. 32694.
- Fowlie, Meredith, Catherine Wolfram, Patrick Baylis, C Anna Spurluck, Annika Todd-Blick, and Peter Cappers. 2021. “Default Effects and Follow-on Behaviour: Evidence from an Electricity Pricing Program.” *Review of Economic*

- Studies* 88 (6):2886–2934.
- Goli, Ali, Jason Huang, David Reiley, and Nickolai M. Riabov. 2025. “Measuring Consumer Sensitivity to Audio Advertising: A Long-Run Field Experiment on Pandora Internet Radio.” *Quantitative Marketing and Economics* :1–31.
- Handel, Benjamin R. 2013. “Adverse Selection and Inertia in Health Insurance Markets: When Nudging Hurts.” *American Economic Review* 103 (7):2643–2682.
- He, Di, Aadharsh Kannan, Tie-Yan Liu, R Preston McAfee, Tao Qin, and Justin M Rao. 2017. “Scale Effects in Web Search.” In *Web and Internet Economics: 13th International Conference, WINE 2017, Bangalore, India, December 17–20, 2017, Proceedings 13*. Springer, 294–310.
- Heidhues, Paul, Alessandro Bonatti, L Elisa Celis, Gregory S Crawford, David Dinielli, Michael Luca, Tobias Salz, Monika Schnitzer, Fiona M Scott Morton, Michael Sinkinson et al. 2023. “More Competitive Search Through Regulation.” *Yale Journal on Regulation* 40 (3):915.
- Ho, Kate, Joseph Hogan, and Fiona Scott Morton. 2017. “The Impact of Consumer Inattention on Insurer Pricing in the Medicare Part D program.” *RAND Journal of Economics* 48 (4):877–905.
- Hovenkamp, Erik. 2024. “The Competitive Effects of Search Engine Defaults.” Working paper.
- Imbens, Guido W. 2020. “Potential Outcome and Directed Acyclic Graph Approaches to Causality: Relevance for Empirical Practice in Economics.” *Journal of Economic Literature* 58 (4):1129–1179.
- Israel, Mark. 2005. “Services as Experience Goods: An Empirical Examination of Consumer Learning in Automobile Insurance.” *American Economic Review* 95 (5):1444–1463.
- Katz, Justin and Hunt Allcott. 2025. “Digital Media Mergers: Theory and Application to Facebook-Instagram.” Working paper.
- Klein, Tobias J, Madina Kurmangaliyeva, Jens Prüfer, and Patricia Prüfer. 2023. “How Important Are User-Generated Data For Search Result Quality? Experimental Evidence.” TILEC Discussion Paper No. 2022–016.
- Klemperer, Paul. 1987. “The Competitiveness of Markets with Switching Costs.” *RAND Journal of Economics* 18 (1):138–150.
- Lee, Kwok Hao and Leon Musolff. 2023. “Entry into Two-sided Markets Shaped by Platform-guided Search.” Working paper.
- Miller, Klaus M, Navdeep S Sahni, and Avner Strulov-Shlain. 2023. “Sophisticated Consumers with Inertia: Long-term Implications from a Large-scale Field Experiment.” Working paper.
- Nelson, Phillip. 1970. “Information and Consumer Behavior.” *Journal of Political Economy* 78 (2):311–329.
- Ostrovsky, Michael. 2021. “Choice Screen Auctions.” In *Proceedings of the 22nd ACM Conference on Economics and Computation*. 741–742.
- Page, Larry. 2012. “2012 Update from the CEO.” Accessed at <https://web.archive.org/web/20120408081225/investor.google.com/corporate/2012/ceo-letter.html> in November 2024.
- Patterson, Mark R. 2013. “Google and Search Engine Market Power.” *Harvard Journal of Law and Technology* 2013-07-01.
- Pearl, Judea and Dana Mackenzie. 2018. *The Book of Why: the New Science of Cause and Effect*. Basic Books.
- Schaefer, Maximilian and Geza Sapi. 2023. “Complementarities in Learning from Data: Insights from General Search.” *Information Economics and Policy* 65 (101063).
- Schmalensee, Richard. 1982. “Product Differentiation Advantages of Pioneering Brands.” *American Economic Review* 72 (3):349–365.
- SEO Statistics. 2024. “124 SEO Statistics for 2024.” Webpage. Accessed at <https://ahrefs.com/blog/seo-statistics/> in March 2025.
- StatCounter. 2024a. “Desktop vs. Mobile vs. Tablet Market Share: United States of America.” Webpage. Accessed at <https://gs.statcounter.com/platform-market-share/desktop-mobile-tablet/united-states-of-america/2023> in April 2024.

- . 2024b. “Search Engine Market Share Worldwide.” Webpage. Accessed at <https://gs.statcounter.com/search-engine-market-share> in April 2024.
- Stigler Committee on Digital Platforms. 2019. “Stigler Committee on Digital Platforms: Final Report.” *Stigler Center News*. Accessed at <https://www.chicagobooth.edu/research/stigler/news-and-media/committee-on-digital-platforms-final-report> in March 2025.
- Tirole, Jean and Michele Bisceglia. 2024. “Fair Gatekeeping in Digital Ecosystems.” TSE Working Paper.
- Tucker, Catherine. 2019. “Digital Data, Platforms and the Usual [Antitrust] Suspects: Network Effects, Switching Costs, Essential Facility.” *Review of Industrial Organization* 54 (4):683–694.
- UK Competition and Markets Authority. 2020. “Online Platforms and Digital Advertising Market Study.” Accessed at <https://www.gov.uk/cma-cases/online-platforms-and-digital-advertising-market-study> in April 2024.
- US House Judiciary Subcommittee on Antitrust. 2020. “Investigation of Competition in Digital Markets.” Accessed at https://democrats-judiciary.house.gov/uploadedfiles/competition_in_digital_markets.pdf in March 2025.
- Varian, Hal. 2015. “Is there a Data Barrier to Entry?” Accessed at <https://www.learconference2015.com/wp-content/uploads/2014/11/Varian-slides.pdf> in March 2025.
- Varian, Hal R. 2007. “Position Auctions.” *International Journal of Industrial Organization* 25 (6):1163–1178.
- Vásquez Duque, Omar. 2022. “The Potential Anticompetitive Stickiness of Default Applications: Addressing Consumer Inertia with Randomization.” Working paper.
- Wernerfelt, Nils, Anna Tuchman, Bradley T. Shapiro, and Robert Moakler. 2024. “Estimating the Value of Offsite Tracking Data to Advertisers: Evidence from Meta.” *Marketing Science* 4 (2):268–286.

Online Appendix

Sources of Market Power in Web Search: Evidence from a Field Experiment

Hunt Allcott, Juan Camilo Castillo, Matthew Gentzkow, Leon Musolff, and Tobias Salz

Table of Contents

A Experimental Results Appendix	2
A.1 Intent-To-Treat Analysis of Browser Extension Treatments	11
B Demand Model Appendix	12
B.1 Identification Details	12
B.2 Estimation Details	17
C Economies of Scale Appendix	19
C.1 Summary Statistics	19
C.2 Descriptives	19
C.3 Implementation Details for Identification Argument	21
C.4 Taylor Expansion to Address HDFE in NLLS	23
C.5 Effect of Data on Result Relevance	23
C.6 Cross-Query Learning	25
D Counterfactuals Appendix	27
D.1 Direct Effects	29
D.2 Equilibrium Effects	31
D.3 Additional Counterfactual Simulation Results	33

A Experimental Results Appendix

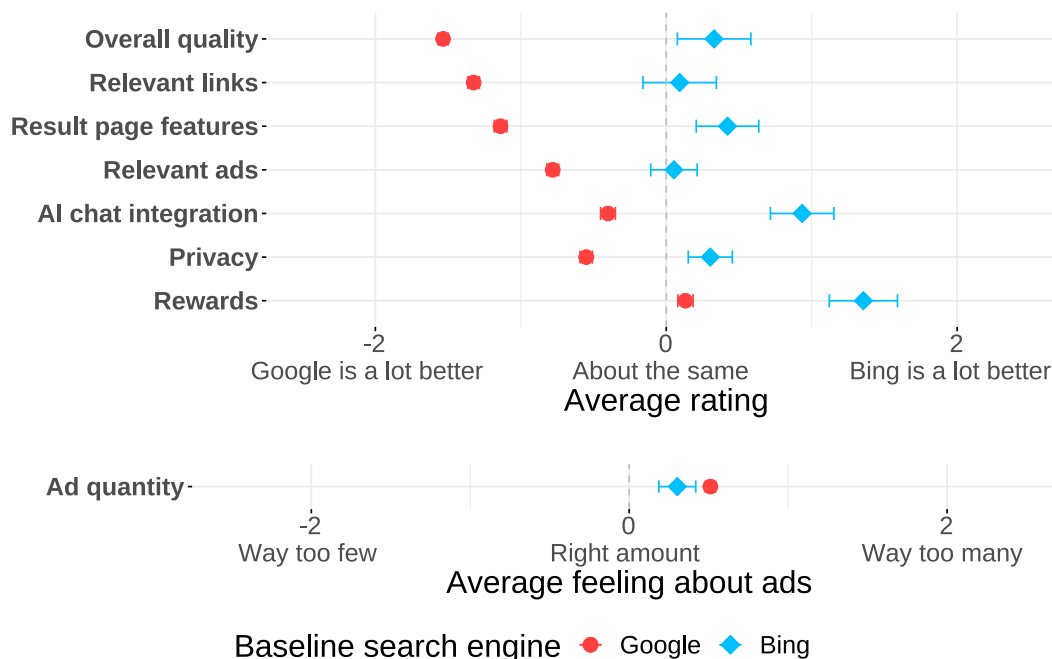


Figure A1: Initial Ratings of Google and Bing (Including Number of Ads)

Notes: This figure presents average responses to the search engine rating questions for baseline Google and Bing users. The top rows present the average rating of Google and Bing on each reported dimension, in response to the following questions: “Overall, how would you rate the quality of Google relative to Bing?” and “How would you rate the quality of Google relative to Bing on the following dimensions?” Response options were “Bing is a lot better,” “Bing is a little better,” “They are about the same,” “Google is a little better,” and “Google is a lot better,” coded as 2, 1, 0, -1, and -2, respectively. The bottom row presents the average response to the following question: “How do you feel about the number of ads on [baseline search engine used]?” Response options were “way too many,” “too many,” “right amount,” “too few,” and “way too few,” coded as 2, 1, 0, -1, and -2, respectively. Whiskers indicate 95 percent confidence intervals.

Table A1: **Covariate Balance**

	(1) Control	(2) Active Choice	(3) Default Change	(4) Switch (\$1)	(5) Switch (\$25)	(6) Switch (\$10 & CC)	(7) Switch (\$10 & BC)	(8) Switch (\$10 & CR)	(9) Switch (\$10 & BR)	F-test p-value
Income (\$000s)	52.30	62.48	59.64	59.07	61.94	53.80	52.18	55.96	57.59	0.40
College Degree	0.46	0.57	0.63	0.57	0.57	0.57	0.59	0.58	0.61	0.58
Male	0.33	0.44	0.52	0.47	0.46	0.42	0.39	0.47	0.48	0.13
Age	35.89	35.64	37.22	34.52	35.89	35.48	36.63	36.80	37.49	0.43
White	0.68	0.67	0.73	0.62	0.69	0.62	0.65	0.68	0.66	0.52

(a) **Baseline Google Users**

	(1) Active Choice	(2) Switch (\$10 & CC)	(3) F-test p-value
Income (\$000s)	45.00	54.21	0.43
College Degree	0.47	0.45	0.82
Male	0.53	0.66	0.25
Age	35.13	37.61	0.37
White	0.66	0.63	0.81

(b) **Baseline Bing Users**

Notes: Panels (a) and (b) present balance tests within the baseline Google user and baseline Bing user samples, respectively, for the variables specified by the row labels. Columns 1–9 in Panel (a) and 1–2 in Panel (b) present means for each treatment group. The rightmost column presents the p-value of an F-test of a participant-level regression of each variable on the treatment group indicators. The sample underlying this table includes all participants (including participants who did not stay with us until endline).

Table A2: Completion Rates

	(1) Control	(2) Active Choice	(3) Default Change	(4) Switch Bonus (\$1)	(5) Switch Bonus (\$25)	(6) Switch Bonus (\$10 & CC)	(7) Switch Bonus (\$10 & BC)	(8) Switch Bonus (\$10 & CR)	(9) Switch Bonus (\$10 & BR)	(10) F-test p-value
Finished Survey 2	0.937	0.994	0.968	0.977	0.92	0.962	0.943	0.935	0.979	0.012
Kept Search Extension 2 weeks after Survey 2	0.873	0.964	0.943	0.93	0.857	0.92	0.874	0.906	0.936	0.009
Kept Search Extension 2 months after Survey 1	0.825	0.898	0.861	0.86	0.804	0.847	0.829	0.834	0.851	0.6

(a) Baseline Google Users

	(1) Active Choice	(2) Switch Bonus (\$10 & CC)	(3) F-test p-value
Finished Survey 2	0.921	0.921	1
Kept Search Extension 2 weeks after Survey 2	0.921	0.868	0.461
Kept Search Extension 2 months after Survey 1	0.842	0.868	0.748

(b) Baseline Bing Users

Notes: Panels (a) and (b) present balanced attrition tests within the baseline Google user and baseline Bing user samples, respectively. Columns 1–9 in Panel (a) and 1–2 in Panel (b) present completion rates for each treatment group. The first row presents the share of participants that completed Survey 2. The second row presents the share of participants that kept Search Extension installed for 14 days after completing Survey 2. The final row presents the share of participants that kept Search Extension installed eight weeks after completing Survey 1. The sample in each row is a strict subset of the row above. The rightmost column presents the p-value of an F-test of a participant-level regression of completion indicators on the treatment group indicators.

Table A3: Average Baseline Ratings of Google and Bing by Completion Status (Google Users)

	(1) Control	(2) Active Choice	(3) Default Change	(4) Switch Bonus (\$1)	(5) Switch Bonus (\$25)	(6) Switch Bonus (\$10 & CC)	(7) Switch Bonus (\$10 & BC)	(8) Switch Bonus (\$10 & CR)	(9) Switch Bonus (\$10 & BR)	(10) F-test p-value
All users	-1.68	-1.48	-1.56	-1.67	-1.52	-1.5	-1.59	-1.51	-1.48	0.204
Finished Survey 2	-1.69	-1.48	-1.56	-1.68	-1.52	-1.5	-1.59	-1.5	-1.48	0.165
Kept Search Extension 2 weeks after Survey 2	-1.73	-1.48	-1.56	-1.68	-1.52	-1.5	-1.57	-1.51	-1.49	0.205
Kept Search Extension 2 months after Survey 1	-1.71	-1.48	-1.57	-1.7	-1.5	-1.48	-1.59	-1.5	-1.49	0.113
Attriters vs Stayers	0.17	-0.05	0.07	0.20	-0.09	-0.18	-0.01	-0.04	0.09	0.87

Notes: The table present balance tests on the baseline (Survey 1) ratings of Bing relative to Google (those presented in Figure 2) among baseline Google users. For the first four rows, columns 1–9 present the average ratings of search engines for each treatment group, and the last column shows the p-values of the F-test associated with a participant-level regression of the baseline rating on the treatment group indicators. The first row presents the average baseline ratings for all users. The second row presents the average baseline ratings of the participants that completed Survey 2. The third row presents the average baseline rating of the participants that kept Search Extension installed for 14 days after completing Survey 2. The fourth row presents the average baseline ratings of the participants that kept Search Extension installed eight weeks after completing Survey 1. For the first four rows, the sample in each row is a strict subset of the row above. The last row presents coefficients that measure differences between attriters and stayers, where stayers are defined as in the fourth row. Specifically, columns 1–9 present the coefficients γ_k from the regression $Y_i = \beta_0 + \sum_{k \in \{\text{Treatments}\}} \beta_k T_{ik} + \sum_{k \in \{\text{Treatments}\}} \gamma_k (T_{ik} \times \text{Attriter}_i) + \varepsilon_i$, where T_{ik} is a treatment dummy for the treatment k for the user i and Attriter_i is an indicator of whether the user i kept the extension installed for two months after Survey 1. The rightmost column presents the F-test p-value for $\gamma_k = 0$ for $k \in \{\text{Treatments}\}$.

Table A4: Search Volume: Searches Per Day

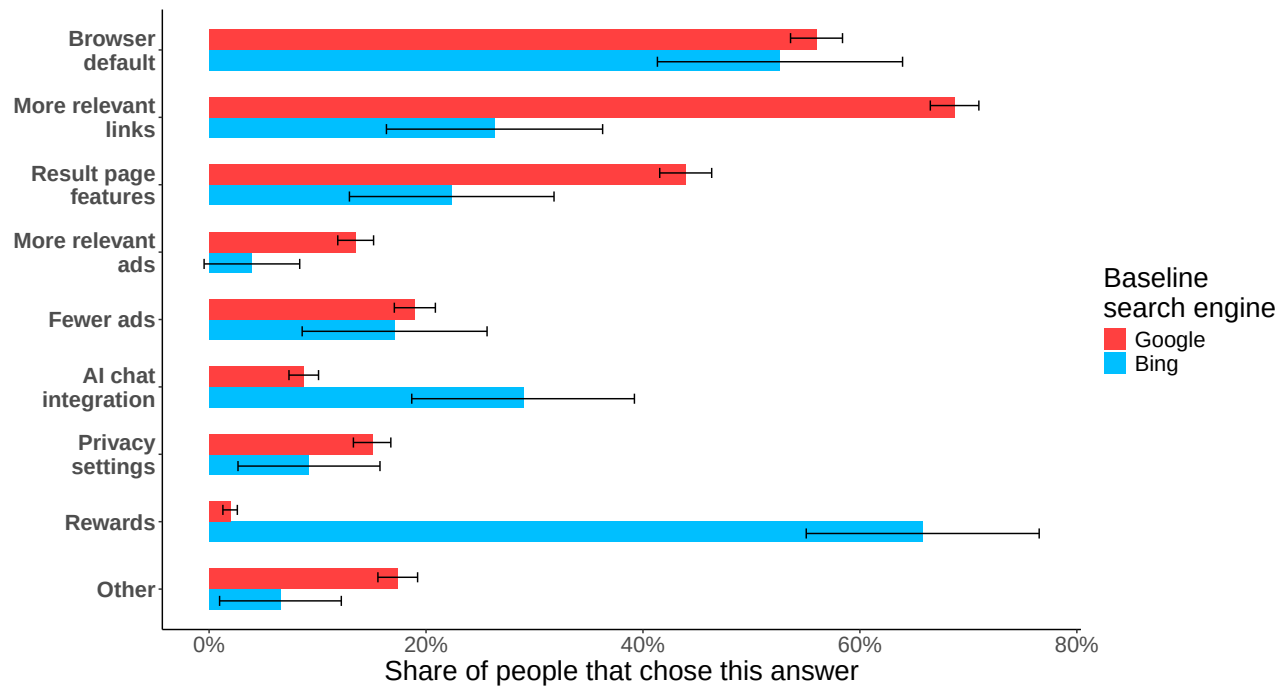
	(1) Control	(2) Active Choice	(3) Default Change	(4) Switch (\$1)	(5) Switch (\$25)	(6) Switch (\$10 & CC)	(7) Switch (\$10 & BC)	(8) Switch (\$10 & CR)	(9) Switch (\$10 & BR)	F-test p-value
t = 0	11.041	14.453	13.002	11.165	12.791	12.707	12.874	13.047	13.336	0.853
t = 1	10.036	14.220	12.564	12.995	15.157	16.194	14.463	14.177	14.905	0.141
t = 2	10.645	13.472	13.685	12.444	13.102	11.959	12.615	11.755	12.625	0.791
p-value (t = 0, t = 1)	0.402	0.665	0.708	0.308	0.089	0.000	0.037	0.085	0.018	-
p-value (t = 0, t = 2)	0.724	0.130	0.223	0.511	0.998	0.104	0.506	0.066	0.213	-

(a) Baseline Google Users

	(1) Active Choice	(2) Switch (\$10 & CC)	F-test p-value
t = 0	15.053	16.855	0.678
t = 1	14.278	14.872	0.872
t = 2	13.689	14.337	0.863
p-value (t = 0 vs t = 1)	0.576	0.258	-
p-value (t = 0 vs t = 2)	0.251	0.192	-

(b) Baseline Bing Users

Notes: This figure presents the average number of searches per day across all search engines, by treatment group (column) and by phase of the experiment (row). Panel (a) presents results for baseline Google users, and Panel (b) presents results for baseline Bing users. The phases are defined as follows: $t = 0$, $t = 1$, and $t = 2$ refer to the the days before Survey 1, the days between Survey 1 and Survey 2, and the days after Survey 2, respectively. The bottom two rows present the p-value of paired t-tests between period 0 and each of the two other periods. The rightmost column presents the p-value of an F-test associated with a participant-level regression of the daily average search volume on the treatment group indicators.

Figure A2: **Why People Use Google or Bing**

Notes: This figure presents the share of baseline Google and Bing users that chose each answer to the following question: “Why do you use [baseline search engine used] instead of [other search engine] for your searches on this web browser? Choose all that apply.” Whiskers indicate 95 percent confidence intervals.

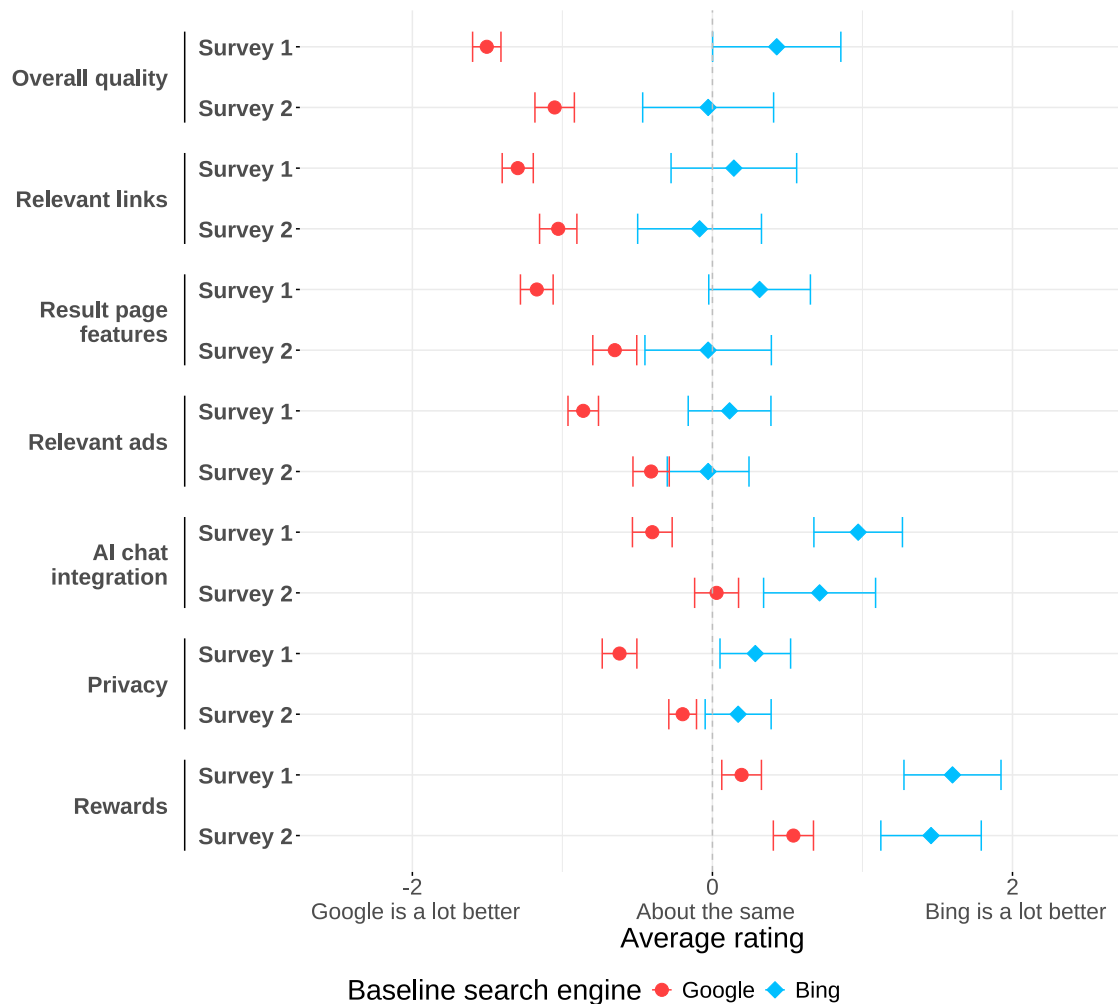


Figure A3: **Switch Treatment (\$10) Search Engine Ratings in Both Surveys**

Notes: This figure presents average responses to the search engine rating questions in each survey. We focus on users in the \$10 Switch Bonus Control group (S10CC), who were asked rating questions in both surveys. The figure presents the average ratings of Google and Bing on each reported dimension, in response to the following questions: “Overall, how would you rate the quality of Google relative to Bing?” and “How would you rate the quality of Google relative to Bing on the following dimensions?” Response options were “Bing is a lot better,” “Bing is a little better,” “They are about the same,” “Google is a little better,” and “Google is a lot better,” coded as 2, 1, 0, -1, and -2, respectively. Whiskers indicate 95 percent confidence intervals.

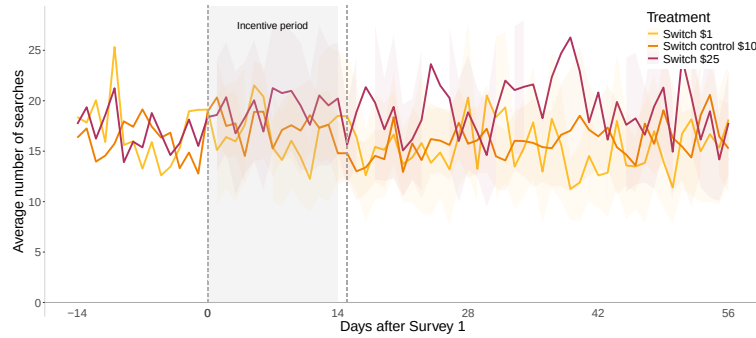


Figure A4: **Search Volume in Switch Bonus Groups**

Notes: This figure presents the average number of searches across all search engines for each day of the experiment. We focus on participants in the Switch Bonus group. The dashed vertical lines mark the dates of the two surveys.

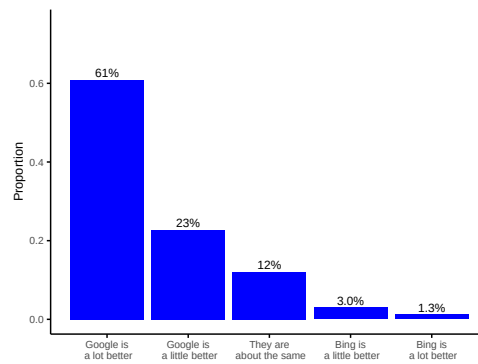


Figure A5: **Baseline Ratings of Google vs. Bing**

Notes: This figure presents the distribution of overall quality ratings reported on Survey 1. The survey question was, “Overall, how would you rate the quality of Google relative to Bing?”

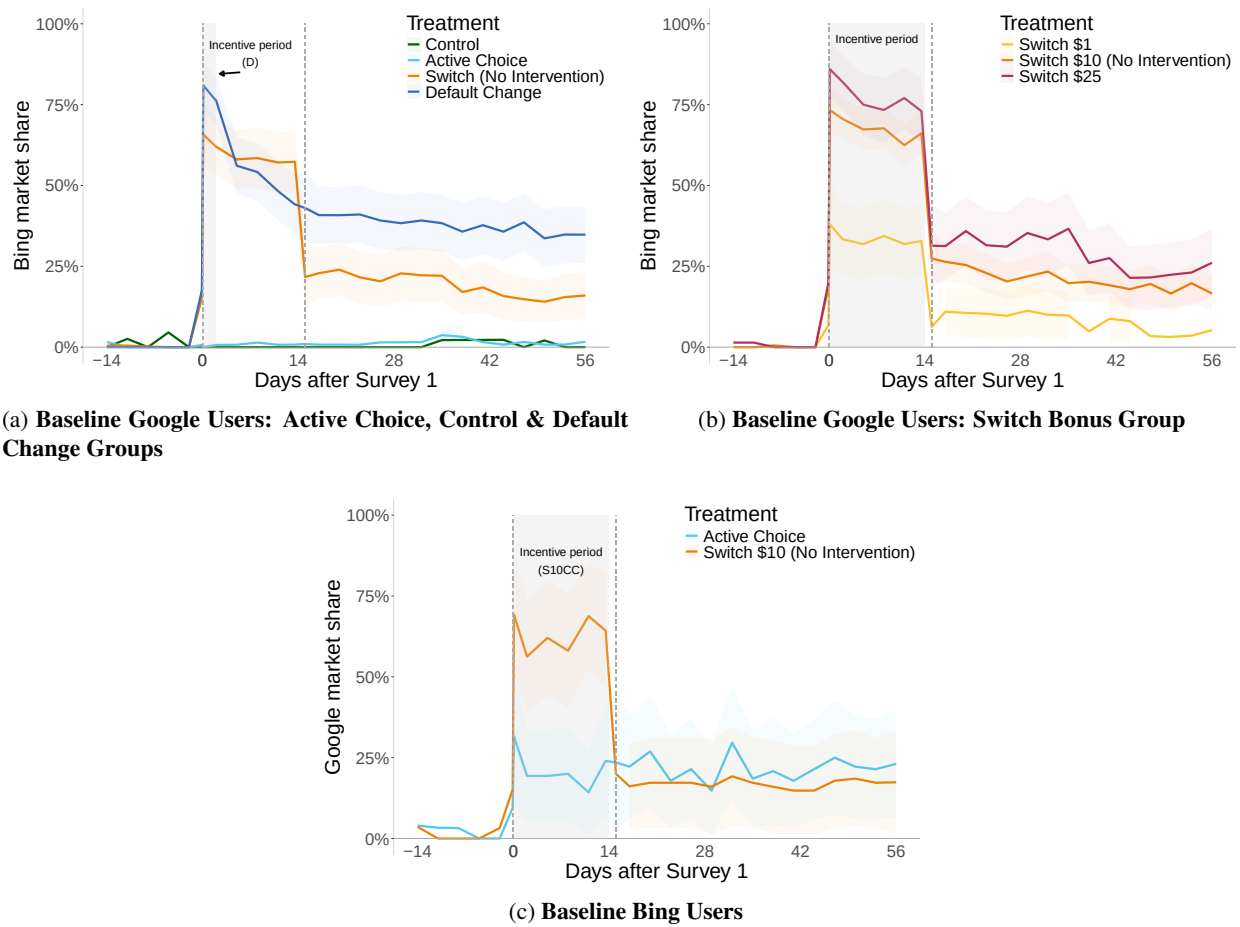


Figure A6: Search Market Shares (Binary) By Treatment Group

Notes: This figure presents the binarized version of Figure 3: the figures are equivalent except that here, to compute daily market shares, we first compute the daily binary choice indicator of non-default search engine for each participant and then calculate the average shares across participants.

A.1 Intent-To-Treat Analysis of Browser Extension Treatments

Table A5: Effects of Ranking Degradation and Ad Blocking on Quality Ratings: Intent-to-Treat

Dep. var:	Panel A: \$10 Switch Group: Change in Ratings of Bing Relative to Google							
	(1) Overall quality	(2) Relevant links	(3) Result page features	(4) Relevant ads	(5) AI chat	(6) Privacy	(7) Rewards	(8) Number of ads
Ad Blocking	-0.038 (0.070)	-0.050 (0.072)	-0.152* (0.086)	-0.169** (0.072)	-0.087 (0.076)	-0.071 (0.061)	-0.118* (0.065)	0.090 (0.060)
Ranking Degradation	-0.287*** (0.070)	-0.311*** (0.072)	-0.328*** (0.086)	-0.165** (0.073)	-0.037 (0.076)	-0.099 (0.061)	0.065 (0.065)	-0.019 (0.060)
Constant	0.425*** (0.063)	0.241*** (0.066)	0.583*** (0.077)	0.432*** (0.068)	0.479*** (0.067)	0.425*** (0.054)	0.365*** (0.057)	-0.050 (0.051)
R ²	0.019	0.021	0.020	0.012	0.002	0.005	0.005	0.003
N	895	895	895	895	895	895	895	895

Notes: This table presents an intent-to-treat version of Panel A of Table 4. It only differs from that table in that it also includes participants in the \$10 Switch Bonus group that declined the offer to switch. We do not provide an ITT version of Panel B as participants who do not accept the offer to switch naturally do not interact with Bing; we also did not pre-register this Panel.

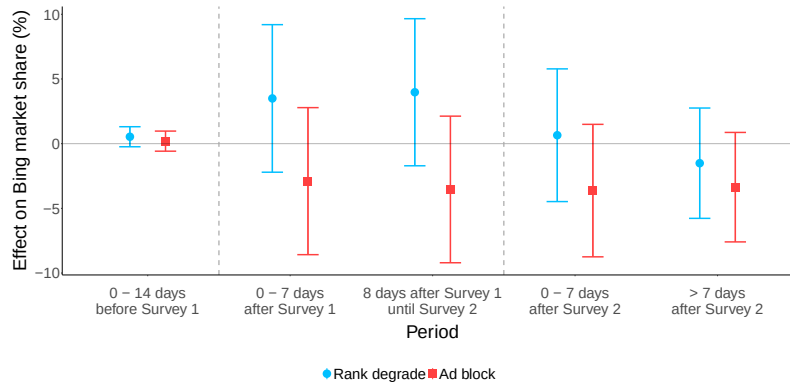


Figure A7: Effects of Ranking Degradation and Ad Blocking on Market Share: Intent-To-Treat

Notes: This figure presents an intent-to-treat version of Table 4. It only differs from that figure in that it also includes participants in the \$10 Switch Bonus group that declined the offer to switch.

B Demand Model Appendix

B.1 Identification Details

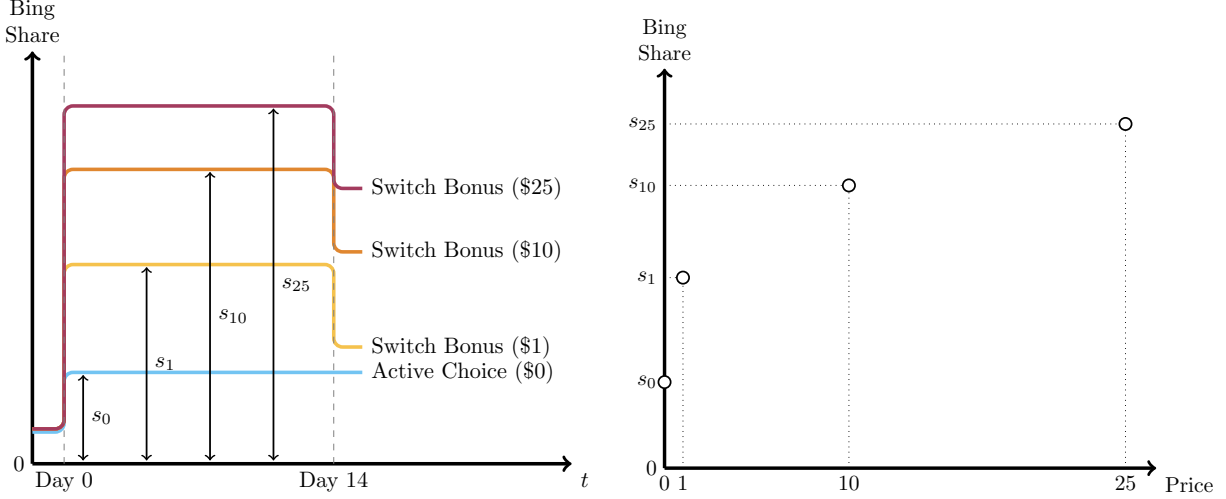


Figure A8: **Identification of Idiosyncratic Preference Distribution**

Notes: This figure illustrates, for Chrome users, the Bing market shares over time in each treatment condition (left). It also illustrates what fraction of users are willing to switch to Bing at any given price (right).

Distribution of idiosyncratic preferences and price response η . By comparing market shares of the Switch Bonus group at different payment offers during the incentivized period, we can identify the distribution of idiosyncratic preferences $S(\cdot)$ and the price response η , as shown in Figure A8. Consider a Switch Bonus user who is offered a price p to use search engine $-d$ after Survey 1. The utility from declining the Switch Bonus and staying with search engine d is

$$\zeta_d^* + \chi_{id} + \delta V_{i,t=2}(d). \quad (13)$$

The utility from accepting the Switch Bonus and switching the default to $j = -d$ is

$$\eta p + \tilde{\zeta}_{-d} + \chi_{i,-d} + \delta V_{i,t=2}(-d). \quad (14)$$

Survey 1 tells participants that, regardless of their $t = 1$ choice, they will be guided through the choice screens on Survey 2 and consumers are therefore forced to pay the switching costs at time $t = 2$, even if they do not switch. Furthermore, we have assumed that consumers believe they know $\tilde{\zeta}_{-d}$ with certainty, so there is no perceived value from exploration. Therefore, $V_{i,t=2}(d) = V_{i,t=2}(-d)$ and the continuation values drop out from the comparison. Having this in mind, the consumer chooses $-d$ if

$$\eta p + \Delta \tilde{\zeta} + \Delta \chi_i > 0. \quad (15)$$

The modeled market share of search engine $-d$ at time $t = 1$ is thus

$$s_{-d,t=1}^{Sp} = S\left(\eta p + \Delta\tilde{\zeta}\right). \quad (16)$$

The parameters η and $\Delta\tilde{\zeta}$ simply play the role of a scale factor and a shifter, so we can rewrite this expression as

$$s_{-d,t=1}^{Sp} = H(p), \quad (17)$$

where $H(\cdot)$ is the cumulative density function of willingness to accept, $-1/\eta \cdot (\Delta\chi_i + \Delta\tilde{\zeta})$, a linear transformation of the idiosyncratic preferences term. Since the left hand side of Equation 17 is data, different price offers directly identify values of $H(\cdot)$ at different points. As is standard in discrete choice models, we normalize the mean and variance of $\Delta\chi_i$. This gives us the distribution $S(\cdot)$ of the normalized error $\Delta\chi_i$ from the shape of $H(p)$.²⁸

When $S(\cdot)$ is known, any two price points identify η , for instance:

$$\eta = \frac{S^{-1}(s_{-d,t=1}^{S25}) - S^{-1}(s_{-d,t=1}^{S1})}{25 - 1}$$

Perceived difference in quality $\Delta\tilde{\zeta}$. The perceived difference in quality is identified by market shares among Active Choice users, as depicted in Figure 5. Since these users made an active choice on Survey 1, we assume that their market shares from that point on are not driven by switching costs or inattention. Instead, they are entirely determined by perceived differences in quality.

Formally, for Active Choice users, the market share of the alternative search engine $-d$ is as in equation (7), with zero price difference Δp and without the switching cost $\sigma(1 - \delta)$:

$$s_{-d,t \geq 1}^A = S(\Delta\tilde{\zeta}). \quad (18)$$

The relevant difference in quality is $\Delta\tilde{\zeta}$, capturing the fact that these users have not used search engine $-d$ and might thus have wrong perceptions about its quality. We can invert this equation to obtain the following expression for the perceived difference in quality $\Delta\tilde{\zeta}$:

$$\Delta\tilde{\zeta} = S^{-1}(s_{-d,t \geq 1}^A), \quad (19)$$

where $S^{-1}(x)$ is the inverse of the cumulative density function of $-\Delta\chi_i$.

Learning $\zeta_{-d}^* - \tilde{\zeta}_{-d}$. To identify learning, we compare the active choices made by Switch Bonus users after they had two weeks to learn about search engine $-d$ with the choices made by Active Choice users, who are not familiar with search engine $-d$ (see Figure 5).

²⁸Technically, $H(\cdot)$ would be non-parametrically identified if we had a switch treatment for each price point p . In that case, $S(\cdot)$ would be non-parametrically identified up to a scale and location normalization. In practice, we have enough price points to determine that a log-normal fits our data well as illustrated in Figure 6: some participants are close to indifferent between Google and Bing whereas others require large payments to abandon Google for two weeks.

When Switch Bonus subjects in the Search Extension Intervention Control group make their Survey 2 active choice, they have had time to learn ζ_{-d}^* , the true quality of search engine $-d$. They thus choose it if

$$\Delta\zeta^* + \Delta\chi_i > 0. \quad (20)$$

Assuming $\eta p_1 > \zeta^* - \tilde{\zeta}$, that is, that the Switch Bonus was large enough that all consumers who would choose Bing under perfect information were induced to try Bing, the market share of $-d$ is thus

$$s_{-d,t \geq 2}^{S10CC} = S(\Delta\zeta^*). \quad (21)$$

Comparing this expression with the market share for Active Choice users (equation 18) and rearranging gives the following expression:

$$\zeta_{-d}^* - \tilde{\zeta}_{-d} = S^{-1}(s_{-d,t \geq 2}^{S10CC}) - S^{-1}(s_{-d,t \geq 1}^A). \quad (22)$$

Attention probability π . To explain how we identify π , we will first explain how the market share in the default group evolves.²⁹ Let $\tilde{s}^*$ be long-run share of the alternative search engine in the D group—i.e., the share after everybody who is not permanently inattentive has made an attentive choice. Also, let s_{-d*}^D be the share of the alternative search engine directly after treatment. Let $\tilde{\pi}$ be the rate at which participants who are not permanently inattentive become attentive in a given week.³⁰ With these definitions, the share of users who still use the alternative search engine in a given week is given by a weighted average of permanently inattentive users and users that become stochastically attentive; the users who become stochastically attentive converge away from s_{-d*}^D to $\tilde{s}^*$.

$$s_{-d, \text{week}=w}^D = \phi s_{-d*}^D + (1 - \phi) [(1 - \tilde{\pi})^w \cdot s_{-d*}^D + (1 - (1 - \tilde{\pi})^w) \cdot \tilde{s}^*]. \quad (23)$$

Intuitively, among those users who (i) are not permanently inattentive and (ii) would like to switch back, only a fraction $\tilde{\pi}$ actually switch back during a given week (which corresponds to half a period in our model). Therefore, the share of users switching away from the alternative search engine in a given week decays geometrically with a rate $\tilde{\pi}$, and we can identify that rate of decay from the following expression (which follows directly from equation (23)):

$$s_{-d, \text{week}=2}^D - s_{-d, \text{week}=1}^D = (1 - \tilde{\pi}) (s_{-d, \text{week}=1}^D - s_{-d*}^D), \quad (24)$$

where s_{-d*}^D is the initial market share directly after treatment—that is, the fraction of users who switched to obtain our payment—and $s_{-d, \text{week}=w}^D$ represents the market share among the D group *at the end* of week w . Hence, we can derive the following expression for π directly as a function of market shares

²⁹Our explanation assumes that learning has not yet occurred, which means that the derived expressions are only correct for week one and two. This suffices for identification of the parameters.

³⁰Given that π is defined for a two-week period, this is given by $\tilde{\pi} = 1 - (1 - \pi)^{1/2}$.

$$\pi = 1 - \left(\frac{s_{-d, \text{week}=2}^D - s_{-d, \text{week}=1}^D}{s_{-d, \text{week}=1}^D - s_{-d, *}^D} \right)^2. \quad (25)$$

Switching costs σ and permanent inattention ϕ . As suggested by Figure 5, the switching cost σ and inattention parameter ϕ are jointly identified from the difference between the Active Choice and Control market share as well as the difference between the long run Default Change and Switch Bonus market share. In particular, switching costs and inattention both create inertia: consumers are less likely to switch from the search engine they previously used, increasing the difference between the Active choice and Control market shares as well as the long-run Default Change market share. However, as we will argue next, σ and ϕ are separately identified because σ affects each of these quantities symmetrically while inattention has a stronger effect on the long-run Default Change market share.

First, both types of inertia create a gap between the shares for Control users—who are subject to both forms of inertia—and for Active Choice users—who are not subject to either. The gap between those market shares is

$$s_{-d, t \geq 0}^A - s_{-d, t \geq 0}^C = S(\Delta \tilde{\zeta}) - (1 - \phi)S(\Delta \tilde{\zeta} - \sigma(1 - \delta)), \quad (26)$$

which is increasing in both σ and ϕ .

Second, both types of inertia lead to higher market shares for the Default Change group after the incentive period. High switching costs and permanent inattention both imply that the geometric decay process that describes the market share over time will settle at a higher level. To state this formally, suppose for this section only that agents learn the true quality of the alternative search engine instantaneously after switching.³¹ We now obtain an expression for $s_{-d, \infty}^D$, the value the market share $s_{-d, t}^D$ converges to as $t \rightarrow \infty$. Let $s^* = S(\Delta \zeta^* + \sigma(1 - \delta))$ be the hypothetical long-run market share of the alternative search engine among D group users, assuming (i) everybody is attentive and (ii) users *have* learned its true quality. With these definitions, we can express the actual long-run market share of the alternative search engine as

$$s_{-d, \infty}^D = \phi s_{-d, *}^D + (1 - \phi)s^* = \phi s_{-d, *}^D + (1 - \phi)S(\Delta \zeta^* + \sigma(1 - \delta)) - S(\Delta \zeta^*), \quad (27)$$

which is indeed increasing in ϕ and σ . To see why $s_{-d, \infty}^D$ is the long-run market share note that everybody who is not permanently inattentive (fraction $1 - \phi$) has made a choice and everybody else (fraction ϕ) is still stuck with the default.³²

We now argue that although both moments (26 and 27) depend on switching costs and inattention, inattention has a much stronger effect on the latter. First, suppose there is no switching cost. In that case, $s_{-d, t \geq 0}^A - s_{-d, t \geq 0}^C = \phi S(\Delta \zeta^*) = \phi s_{-d, t \geq 0}^A$. Then note that permanent inattention affects both expressions as follows:

³¹We make this assumption only in this section to simplify the exposition. We otherwise maintain the assumption that people learn after fourteen days. The intuition extends to that case.

³²We do not observe choices at $t = \infty$ as our sample ends after eight weeks, so our estimation (Section 5.2) uses the market share at the end of our experiment. Given our estimates from Section 5.3, the probability of paying attention after eight weeks is on the order of one thousandth, so the difference between these two expressions is negligible.

$$\frac{\partial}{\partial \phi}(s_{-d,t \geq 0}^A - s_{-d,t \geq 0}^C) = s_{-d,t \geq 0}^A \quad \text{and} \quad \frac{\partial}{\partial \phi}s_{-d\infty}^D = s_{-d*}^D - s^* \quad (28)$$

Thus, it affects the gap between the A and C groups to the extent that a lot of people in the active choice group want to use $-d$. As we saw in Section 4, only about 5.63 percent of Chrome users want to use Bing, so permanent inattention will have little effect on the first expression. On the other hand, ϕ has a large impact on the long-run D share $s_{-d\infty}^D$ as long as (i) our treatment induces a large fraction of people s_{-d*}^D to use $-d$ in return for a payment, and (ii) many would not want to use it without payment, i.e. s^* is small. Conditions (i) and (ii) are both true in the data, as we saw in Section 4: over 75 percent of users switch in response to our payment, while the fraction who actually want to use is around 20 percent. Based on these numbers, we should expect the effect of permanent inattention on $s_{-d\infty}^D$ to be on the order of ten times larger than the effect on $s_{-d,t \geq 0}^A - s_{-d,t \geq 0}^C$.

Now suppose that there is no permanent inattention. Then $s_{-d,t \geq 0}^A - s_{-d,t \geq 0}^C = S(\Delta\zeta^*) - S(\Delta\zeta^* - \sigma(1 - \delta))$ and $s_{-d\infty}^D = s^*$. Switching costs affect both expressions as

$$\frac{\partial}{\partial \sigma(1 - \delta)}(s_{-d,t \geq 0}^A - s_{-d,t \geq 0}^C) = S'(\Delta\zeta^* - \sigma(1 - \delta)) \quad \text{and} \quad \frac{\partial}{\partial \sigma(1 - \delta)}s_{-d\infty}^D = S'(\Delta\zeta^* + \sigma(1 - \delta)). \quad (29)$$

Therefore, the effect on both expressions should be roughly similar as long as the density of marginal users does not change too much.

The main takeaway from this analysis is that switching costs roughly have the same impact on $s_{-d,t \geq 0}^A - s_{-d,t \geq 0}^C$ and $s_{-d\infty}^D$, whereas permanent inattention has a much larger impact on $s_{-d\infty}^D$. This provides an argument why both parameters are separately identified.

Quality preferences α and ρ . Comparing the \$10 Switch Bonus at time $t = 2$ across the Ranking Degradation and Ad Blocking conditions identifies preferences for the components of quality. Similar to equation 21, the market share for a Switch Bonus user in the Search Extension Intervention $I \in \{RC, CA, RA\}$ is given by

$$s_{-d,t \geq 2}^{S10I} = S(\Delta\zeta^I). \quad (30)$$

where $\Delta\zeta^I$ is the quality implied by intervention I . Note that if $I = CC$ —that is, if the user was assigned to the control group for both Ranking Degradation and Ad Blocking—then $\Delta\zeta^I = \Delta\zeta^*$.

Recall that search engine quality is given by $\zeta_j = \alpha a_j + \rho r_j + \xi_j$. The effect of Ranking Degradation on quality difference is $\zeta_{-d}^{RC} - \zeta_{-d}^{CC} = \zeta_{-d}^{RA} - \zeta_{-d}^{CA} = \rho(r_{-d}^{RC} - r_{-d}^{CC})$. Thus, comparing Ranking Degradation relative to its control on $t \geq 2$ market shares (equations 21 and 30) and rearranging gives

$$\rho = \frac{S^{-1}(s_{-d,t \geq 2}^{S10RC}) - S^{-1}(s_{-d,t \geq 2}^{S10CC})}{r_{-d}^{RC} - r_{-d}^{CC}}. \quad (31)$$

A similar expression can be obtained by comparing $s_{-d,t \geq 2}^{S10RA}$ and $s_{-d,t \geq 2}^{S10CA}$. Let r_j be defined in units of click through rates. Then, the right-hand side of equation (31) is observed in the data: it is the ratio of two

Table A6: **Empirical Moments for Demand Estimation**

Description	(1) Formula	(2) Estimate	(3) SE
Baseline Chrome users			
Baseline and Control, $t \geq 0$	$\hat{s}_{-d,t \geq 0}^C$	0.013	0.003
Active Choice group Bing share, $t \geq 1$	$\hat{s}_{-d,t \geq 1}^A$	0.026	0.013
Default Change group Bing share, $t = *$	$\hat{s}_{-d,t=*}^D$	0.771	0.040
Default Change group Bing share, $t = 1$	$\hat{s}_{-d,t=1}^D$	0.470	0.047
Default Change group Bing share, $t = 2$	$\hat{s}_{-d,t=2}^D$	0.397	0.047
\$1 Switch Bonus group Bing share, $t = 1$	$\hat{s}_{-d,t=1}^{S1}$	0.337	0.056
\$25 Switch Bonus group Bing share, $t = 1$	$\hat{s}_{-d,t=1}^{S25}$	0.750	0.046
\$10 Switch Bonus (CC) group Bing share, $t = 1$	$\hat{s}_{-d,t=1}^{S10}$	0.677	0.032
\$10 Switch Bonus (CC) group Bing share, $t \geq 2$	$\hat{s}_{-d,t \geq 2}^{S10}$	0.206	0.027
\$10 Switch Bonus (CR) group Bing share, $t \geq 2$	$s_{-d,t \geq 2}^R$	0.158	0.024
\$10 Switch Bonus (BC) group Bing share, $t \geq 2$	$s_{-d,t \geq 2}^B$	0.146	0.025
\$10 Switch Bonus (BR) group Bing share, $t \geq 2$	$s_{-d,t \geq 2}^{BR}$	0.169	0.024
Baseline Edge users			
Baseline Google share, $t = 0$	$\hat{s}_{-d,t=0}$	0.273	0.104
Active Choice group Google share, $t \geq 1$	$\hat{s}_{-d,t \geq 1}^A$	0.368	0.112
\$10 Switch Bonus (CC) group Google Share, $t = 1$	$\hat{s}_{-d,t=1}^{S10}$	0.709	0.106
\$10 Switch Bonus (CC) group Google Share, $t \geq 2$	$\hat{s}_{-d,t \geq 2}^{S10}$	0.389	0.113

Notes: This table presents the empirical moments used for the demand estimation procedure described in Section 5. Standard errors clustered at the participant level.

treatment effects. Analogous equations also hold for the Ad Blocking condition, where we define a_j such that observed Bing ad load is $a_j = 1$.

B.2 Estimation Details

We now explain in detail the moments we use in our GMM procedure; we exhibit the both the empirical and predicted values of these moments in Table A6. The first set of moments are simply market shares: the baseline market share $s_{-d,0}$, the Active Choice market share $s_{-d,t \geq 1}^A$, the market shares for the Switch Bonus group during the incentivized period at different prices $s_{-d,t=1}^{S1}$, $s_{-d,t=1}^{S10CC}$, and $s_{-d,t=1}^{S25}$, and the post-Survey 2 market shares of the Switch Bonus group under different interventions $s_{-d,t \geq 2}^{S10CC}$, $s_{-d,t \geq 2}^{S10RC}$, $s_{-d,t \geq 2}^{S10CA}$, and $s_{-d,t \geq 2}^{S10RA}$. To write out these nine moment conditions, we use m to index the moments that we target. For example, m can represent baseline choices for Chrome users, S10CC choices during the incentivized period for Edge users, or Active Choice choices at time $t \geq 2$ for Chrome users. We denote by $s_m(\theta)$ the market share predicted by our model for moment m when the model parameters are θ . We also define y_{mi} to be subject i 's choice corresponding to moment m . Our first nine moment conditions take the form

$$g_{mi}(\theta) = y_{mi} - s_m(\theta), \quad \mathbb{E}[g_{mi}(\theta^*)] = 0, \quad (32)$$

where θ^* is the vector of true parameters.

To identify the attention probability π , we exploit the market shares of the Default Change group right after Survey 1, after one week, and after two weeks. Rather than using these three market shares directly, we exploit our expression for the identification of π (equation 24). The moment condition that we use is

$$g_{mi}(\theta) = (1 - \pi)^{1/2}(y_{m,i,\text{week}=1} - y_{m,i,*}) - (y_{m,i,\text{week}=2} - y_{m,i,\text{week}=1}), \quad \mathbb{E}[g_{mi}(\theta^*)] = 0 \quad (33)$$

where $y_{m,i,*}$ denotes D group choices right after survey 1, and $y_{m,i,\text{week}=1}$ and $y_{m,i,\text{week}=2}$ denote D group choices at the end of weeks 1 and 2.

To identify switching costs and inattention, we need moments for $s_{-d,t \geq 0}^A - s_{-d,t \geq 0}^C$ and $s_{-d\infty}^D$. We already included moments corresponding to s_{-d0} (which is the same as $s_{-d,t \geq 0}^C$) and $s_{-d,t \geq 1}^A$, so we must include an additional moment for $s_{-d\infty}^D$. We use an empirical analogue of equation (27),

$$g_{mi}(\theta) = y_{m,i,\infty} - \phi y_{m,i,*} - (1 - \phi) [\tilde{s}^*(\theta) + (1 - \pi)^2 (s^*(\theta) - \tilde{s}^*(\theta))], \quad \mathbb{E}[g_{mi}(\theta^*)] = 0 \quad (34)$$

where $y_{m,i,*}$ denotes D group choices right after survey 1, and $y_{m,i,\infty}$ denotes D group choices after a long period has occurred. In practice, we do not observe choices more than two months after the experiment starts, so our actual estimation uses an adjusted version of this moment that uses D group choices at the end of our experiment.³³ However, given our estimates from Section 5.3, the probability of not having paid attention after eight weeks is on the order of one thousandth, so the difference between these two moments is negligible.

There are two important issues we must deal with before computing these moment conditions and the GMM objective function. First, given the nature of our experiment, we don't observe all moments for every participant. For a participant that was randomized into S10CC, for instance, we observe the moments corresponding to S10CC choices at times $t = 1$ and $t \geq 2$, but we don't observe any of the Default Change or Active Choice choices. Second, we used different randomization probabilities for original Google and Bing users, so unconditional means would overweight Bing users in most of our treatments.

To address these issues, we think of our experiment using a potential outcomes setup. Hypothetically, for every moment m , there is a hypothetical realization of $g_{mi}(\theta^*)$. However, because of randomization, we do not observe many of these choices and thus, cannot compute the corresponding moments. To address this issue, we rewrite our moment conditions in the form

$$\tilde{g}_{mi}(\theta) = w_{mi} \cdot g_{mi}(\theta), \quad \mathbb{E}[\tilde{g}_{mi}(\theta^*)] = 0$$

where w_{mi} are weights that allow us to account for the fact that some of the moments $g_{mi}(\theta)$ are unobserved.

³³We now derive the expression for $s_{-d,\text{week}=8}^D$ that we use for estimation. After two weeks (that is, after learning) the fraction of people that would like to switch if attentive is no longer $s_{-d*}^D - \tilde{s}^*$ but $s_{-d*}^D - s^*$. The geometric decay process thus resembles equation 23, but it goes from s_{-d*}^D to s^* (and not from s_{-d*}^D to $\tilde{s}^*$). After accounting for the fraction $[1 - (1 - \tilde{\pi})^2] (s^* - \tilde{s}^*)$ of users that switched back too early, we obtain the following expression for the market shares after week 2:

$$s_{-d,\text{week}=w>2}^D = \phi s_{-d*}^D + (1 - \phi) \left[s^* + \left(1 - (1 - \tilde{\pi})^2 \right) (s^* - \tilde{s}^*) + (1 - \tilde{\pi})^w (s_{-d*}^D - s^*) \right].$$

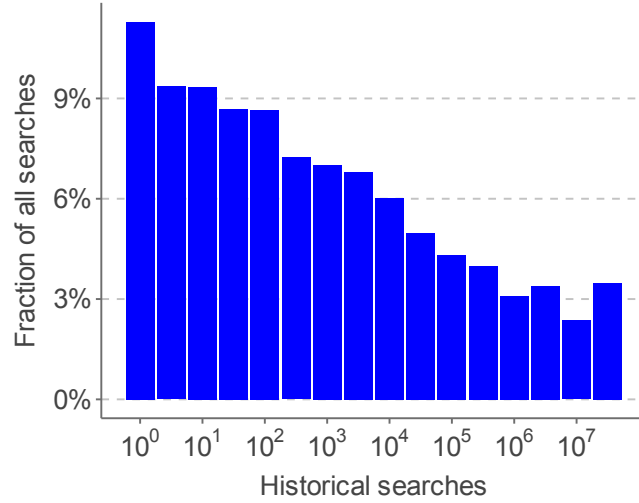


Figure A9: Distribution of Views Across Queries

Notes: This figure provides the fraction of all searches for queries of a certain popularity. We count all searches made on Bing in the period between October 1st, 2021 and October 1st, 2022, and group queries by how often they were searched for. The shape of the distribution supports the commonly-held assumption that there are a large number of tail queries each associated with a lower number of searches but jointly responsible for a considerable share of searches on Bing.

Whenever $g_{mi}(\theta)$ is not observed, we simply set $w_{mi} = 0$ and $\tilde{g}_{mi}(\theta) = 0$. When $g_{mi}(\theta)$ is observed, we set w_{mi} to be the inverse of the (empirical) probability that we observe $g_{mi}(\theta)$ conditional on i 's baseline search engine. Under these weights, it is still the case that $\mathbb{E}[\tilde{g}_{mi}(\theta^*)] = 0$ despite the fact that some of these choices are unobserved and that this occurs with different probabilities for baseline Bing and Google participants.

C Economies of Scale Appendix

C.1 Summary Statistics

Figure A9 presents the distribution of Bing views across all queries between October 1st, 2021 and October 1st, 2022. Table A7 presents summary statistics for the click-and-query data we use to estimate economies of scale in data.

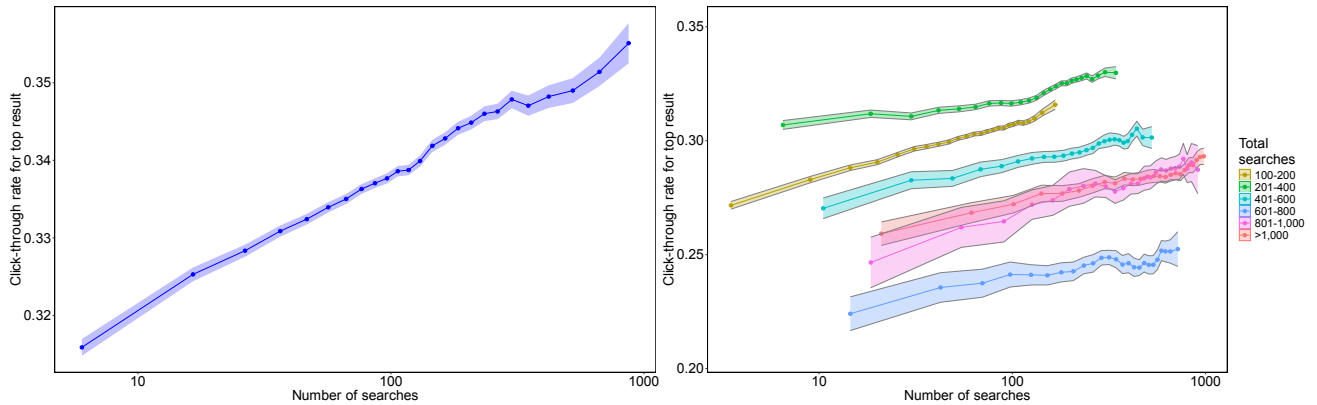
C.2 Descriptives

Before moving on to imposing a functional form, we use binscatters to present nonparametric plots of the relationship between the predicted click-through rate and the number of searches, controlling for query fixed effects. The left plot of Figure A10 exhibits the overall relationship, which seems to be roughly log-linear: each additional doubling in the number of searches leads to an about equal increase in click-through rate. The right plot separately analyzes this relationship for queries of differing popularity: while we find that the average level of click-through rate varies by query popularity, the relationship seems to be robustly well-described as linear in the log of the number of views.

Table A7: **Summary Statistics for Economies of Scale**

Description	Variable	Mean	Min	Median	Max
Number of Searches	n_{qt}	374	1	155	2,980
Dummy: First Result Clicked?	r_{qt}	0.23	0	0	1
Predicted Dummy: First Result Clicked?	$\hat{r}_{qt}$	0.31	0.02	0.18	1.14
Number of Observations		12,194,034			
Number of Queries		43,991			
Number of Top-Ranked URLs		244,136			

Notes: This table provides summary statistics for the dataset underlying the economies of scale analysis.

Figure A10: **Non-Parametric Estimates of Returns to Scale**

Notes: These graphs show the over-time relationship between the number of searches for a particular query and the rate of clicks on the top result (as disciplined by changes in the result shown first.) Formally speaking, we follow Section 6.1 and first regress the click-through rate on fixed effects for the top-ranked URL. We then use binsreg to analyze the relationship between the predicted click-through rate and the number of searches, controlling for query FE. The left plot exhibits the overall relationship, which seems to be roughly log-linear: each additional doubling in the number of views leads to an about equal increase in click-through rate. The right plot separately analyzes this relationship for queries of differing popularity: while we find that the average level of click-through rate varies by query popularity, the relationship seems to be robustly well-described as linear in the log of the number of views.

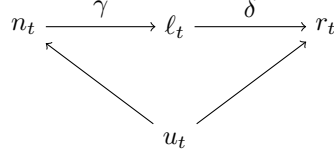


Figure A11: **Directed Acyclic Graph Underlying Estimation Strategy**

Notes: This DAG represents how we identify the causal effect of increasing the number of searches (n_t) on the result relevance as measured by CTR (r_t) while accounting for an unobserved confounder (u_t) that could (e.g.) represent how different types of users arrive over the lifecycle of a query. The key to our identification strategy is the mediator ℓ_t , which represents the search engine’s ranking of links. In particular, the search engine only learns from user data in the form of additional searches, and hence the relationship between n_t and ℓ_t is not confounded. Similarly, the relationship between ℓ_t and r_t is not confounded conditional on n_t .

C.3 Implementation Details for Identification Argument

Our identifying assumption — that the causal effect of more data can only operate through the search ranking — allows us to apply the front door criterion (Pearl and Mackenzie 2018; Imbens 2020; Bellemare, Bloem, and Wexler 2024) to purge confounding variation in click-through rates, such as the observed secular trends. Focusing for now on a single query, the directed acyclical graph (DAG) in Figure A11 illustrates our strategy. In this graph, n_t refers to the number of searches for a given query by time t , ℓ_t to the the rankings of links on the results page at t , and r_t to the click-through rate (i.e., our measure of result relevance). Finally, u_t is an unobserved confounder that affects both the number of prior searches and the click-through rate, such as a changing composition of users over the lifetime of a query. For now, we assume these are all scalars (though they will not be in our eventual implementation.) The causal effect of additional prior searches (i.e., additional data) on the links served is given by γ and the causal effect of the links served on click-through rate is δ .

Regressing r_t directly on n_t would be biased by the confounding variation in u_t . The key insight from this graph is that by regressing r_t on ℓ_t after conditioning on n_t , one can isolate the causal variation in data that leads to changes in click-through rates. The intuition is that we stack many event studies like Figure 7: like in the event-study, we isolate the effect on CTR that comes with a change in the ranking. We then regress the predicted value $\hat{r}_t = \hat{\delta} \cdot \ell_t$ on n_t and obtain the causal effect as the product $\gamma \cdot \delta$.

More formally, the identification challenge is that u_t may introduce a correlation between r_t and n_t , since $u_t \rightarrow n_t$ and $u_t \rightarrow r_t$. The regression

$$r_t = \alpha + \beta n_t + \varepsilon_t$$

would therefore lead to biased estimates of β , i.e. $\lim_{n \rightarrow \infty} \hat{\beta} \neq \gamma \times \delta$.

However, the confounder does not affect the ranking quality directly: there is no arrow from u_t to ℓ_t . Hence, the following regression (Bellemare, Bloem, and Wexler 2024, eq. 7)

$$\ell_t = \kappa + \gamma n_t + \omega_t$$

will yield an unbiased estimator $\hat{\gamma}$ of the effect of additional data (n_t) on ranking of links (ℓ_t). Similarly, we can run (Bellemare, Bloem, and Wexler 2024, eq. 8)

$$r_t = \lambda + \delta \ell_t + \phi n_t + v_t$$

to get an unbiased estimator $\hat{\delta}$ of the effect of the ranking of links (ℓ_t) on CTR (r_t). Multiplying together these two numbers we get an unbiased estimator $\hat{\gamma} \times \hat{\delta}$ of the effect of additional data on CTR.

In practice, we do not literally follow this recipe because of the complication introduced by the fact that the ranking of links ℓ_t is not a scalar. Keeping with it being a scalar, we now explain our alternative recipe. We predict CTR from ranking quality, i.e., we run

$$r_t = \lambda + \delta \ell_t + \phi n_t + v_t$$

and form a prediction $\hat{r}_t = \hat{\delta} \ell_t$ of CTR based just on the current ranking quality. Then we regress this prediction on the number of prior searches n_t , i.e. we effectively run

$$\hat{r}_t = \psi + \eta n_t + \varepsilon_t.$$

As $\hat{r}_t$ here is just ℓ_t multiplied by $\hat{\delta}$, this regression must yield $\hat{\eta} = \hat{\delta} \times \hat{\gamma}$, i.e., our estimator is numerically equivalent to what Bellemare and Bloem's recipe would find.

Now, we introduce our complication: ranking quality is measured by the identity of the top-ranked URL. Let $u(t)$ index the URL top-ranked as of search t (recall we still assume there is just one search term, so we do not need indices for search terms.) Then effectively ℓ_t is a vector of dummies: assuming there are U possible URLs that could be ranked first for this query,

$$\ell_t = (1(u(t) = 1), \dots, 1(u(t) = U))'.$$

This multidimensionality of ℓ_t makes the regression that Bellemare, Bloem, and Wexler (2024) propose hard to interpret and implement.

Still, our alternative way of first forming predictions of CTR works, and as argued above, in the scalar setting it would be exactly equivalent to employing the front-door criterion. Intuitively, we first project CTR on a fixed effect for the top-ranked URL while flexibly controlling for the number of searches a query has received so far. Subsequently, we use the fitted estimates from just the query-by-URL fixed effect in this regression as our dependent variable in estimating the relationship between searches and click-through rate. More formally, indexing queries by q and time by t , we first project CTR on a fixed effect for the top-ranked URL while controlling for a fixed effect for the number of searches a query has received so far, i.e.,

$$r_{qt} = \delta_{q,u(q,t)} + \eta_{n(q,t)} + \varepsilon_{qt}, \quad (35)$$

where $u(q,t)$ gives the index of the top-ranked result served on the search result page and $n(q,t)$ gives the number of searches that query q has seen by time t . As the regression includes a fixed effect $\eta_{n(q,t)}$,

we are flexibly controlling for the number of searches a query has received so far. We then use the fitted estimates $\hat{r}_{qt} = \delta_{q,u(q,t)}$ from just the query-by-URL fixed effect as our dependent variable in estimating the relationship between searches and click-through rate. Intuitively, these fitted estimates will capture systematic improvements in CTR that are driven by Bing changing the order in which it serves search results; by contrast, they will ignore changes due to pure temporal patterns (such as a secular trend in CTR.)

C.4 Taylor Expansion to Address HDFE in NLLS

Our main estimating equation (36) describes a non-linear relationship between the number of previous searches for a query and its click-through rate. However, the estimating equation also features a high-dimension fixed-effect, which is computationally challenging to estimate. To address this concern, this appendix develops a methodology that utilizes repeated Taylor expansions of an estimating equation to derive exact estimates of non-linear parameters in the presence of fixed-effects.

To begin with, we find initial estimates $(\hat{\beta}^0, \hat{\theta}^0)$ by regressing demeaned c_{qt} on demeaned $\frac{\beta}{1-\theta} n_{qt}^{1-\theta}$ (computing this term and then demeaning for any trial value of the parameters.) However, as the underlying regression is not linear, the resulting estimates from this exercise are possibly poor approximations to the true parameter values. To make progress, we turn the regression into a linear problem by utilizing a Taylor series expansion of 9 around initial estimates $(\hat{\beta}^0, \hat{\theta}^0)$. In particular, letting $\hat{\gamma}^0 = \frac{\hat{\beta}^0}{1-\hat{\theta}^0}$, we have

$$\hat{r}_{qt} = \alpha_q + \gamma n_{qt}^{1-\hat{\theta}^0} + \gamma(\hat{\theta}^0 - \theta) \log(n_{qt}) n_{qt}^{1-\hat{\theta}^0} + O(\theta - \hat{\theta}^0)^2 + O(\gamma - \hat{\gamma}^0)^2 + \varepsilon_{qt}.$$

As this equation is linear in easily constructed regressors $n_{qt}^{1-\hat{\theta}^0}$ and $\log(n_{qt}) n_{qt}^{1-\hat{\theta}^0}$, it can be estimated while correctly accounting for the FE α_q , thus yielding new estimates $(\hat{\beta}^1, \hat{\theta}^1)$. We can then form a new Taylor expansion around those estimates, yielding $(\hat{\beta}^2, \hat{\theta}^2)$ and so on. We iterate until convergence, and obtain standard errors via block-boostrapping (resampling at the query-level.)

C.5 Effect of Data on Result Relevance

We now use our estimates in Table 6 to anticipate by how much Bing's CTR would increase if it were to obtain additional data, which could come either from an increase in its market share or from regulatory provisions that require the sharing of click and query data. Note that all estimates in this subsection take a partial equilibrium approach, i.e., they do not consider the effect that an improvement in Bing's CTR may have on its market share and the feedback loop that could potentially result from this effect. When moving to the full model below, we will take into account this flywheel.

Suppose Bing was to obtain an additional 1,000 searches on each query. This would result in an increase of CTR from 23.5 percent to 25.0 percent, an increase of 1.55 percentage points. We can see in Figure A12 that this increase mostly comes from an improvement in serving results on tail queries. Similarly, what if Bing multiplied its market share by 4.28, making it roughly equal to Google's market share? In this case, our estimates imply that Bing's CTR would increase from 23.5 percent to 24.8 percent, an increase of 1.29 percentage points.

More generally, we can use our parameter estimates to calculate counterfactual average click-through

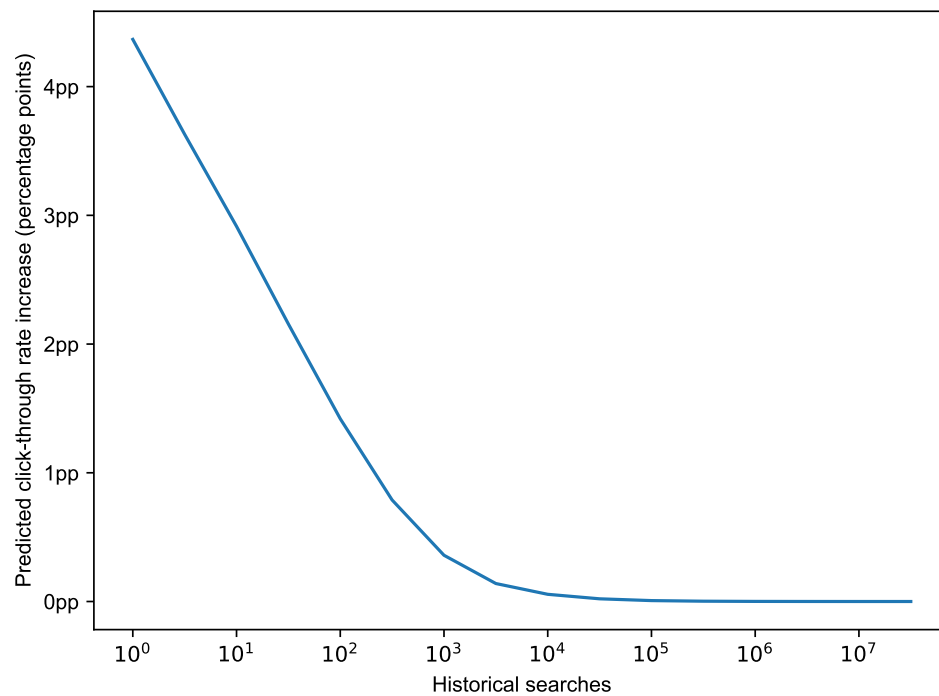


Figure A12: **Marginal Effect of Additional Impressions on Bing CTR by Query Popularity**

Notes: This figure shows the effect (in percentage points) on CTR of increasing the number of searches that Bing observes for each query by 1,000. The resulting improvement in CTR is concentrated on tail queries, which benefit the most from additional data; queries that already had large quantities of data improve less.

rates on Bing if its market share were to be multiplied by λ : the new click-through rate c' after such an increase in market share is given as a function of the old click-through rate by

$$r' = \alpha - \frac{\beta}{1 - \theta} + \lambda^{1-\theta} \cdot \left(r - \alpha + \frac{\beta}{1 - \theta} \right)$$

C.6 Cross-Query Learning

A complication that we do not address in the main text is that learning how to rank results on query q may not be limited to using data from query q —customer behavior on impressions on related queries q' are also helpful. Denote the set of all queries (including the focal query q) as Q , and define a distance metric $d(q, q')$ that measures the distance between any focal query q and potentially related query q' . The distance between any query and itself is zero, i.e., $d(q, q) = 0$ for all q . Similar to in the main text, we will assume that the CTR on query q is given

$$\hat{r}_{qt} = \alpha_q + \beta \frac{1}{1 - \theta} \left(\sum_{q' \in Q} f(d(q, q'); \gamma) \times n_{q't} \right)^{1-\theta} + \varepsilon_{qt} \quad (36)$$

As before, α_q is a fixed effect that captures the fact that different queries may have different baseline CTR. More importantly, β is our measure of the value of gathering additional data, and θ measures economies of scale, i.e., the speed at which this value declines with the amount of data already gathered. In particular, $\theta \approx 0$ implies linear returns to additional data, $\theta \approx 1$ implies logarithmic returns, and $\theta > 1$ implies worse-than-logarithmic returns.

Finally, $f(\cdot)$ is a function (parameterized by γ) that maps the distance $d(q, q')$ between a focal query q and a related query q' into a monotonically declining weight. The speed at which these weights decay as we consider more and more distant queries play an important role in the economies of scale: if the weight decays only slowly, it may not matter if a search engine has never seen a query before as it can apply its learnings from other, related queries. If the weight decays quickly, on the other hand, then not having seen a particular query before would be a strong disadvantage in trying to serve its results. Given our limited data, we will parameterize

$$f(d) = \exp(-\exp(\gamma)d), \quad (37)$$

so that $\gamma = -\infty$ corresponds to no decay with distance (i.e., all views on all queries matter to CTR on any focal query) and $\gamma \rightarrow \infty$ corresponds to the case of no cross-query learning (but still allows views on the focal query to matter as $\lim_{\gamma \rightarrow \infty} \exp(-\exp(\gamma) \times 0) = 1$.)

To estimate (36), we supplement the data on new queries by obtaining, for each query in the original dataset, similar data on the 50 other queries most related to the original query, as reported by Bing's internal metrics. We emphasize that this means we have data only on the most related queries in Q ; to the extent that there is little cross-query learning, we would expect this to not bias our results as searches for less related queries would not yield additional learning on Bing's side. The supplemental dataset contains, aggregated to the daily level, the total number of searches for and clicks on results pages of each of these related queries,

again between 2022-01-24 and 2023-01-23. We note that, by construction, these related queries are not necessarily new; hence, we also obtain the total number of impressions between 2021-01-24 and 2022-01-23 for the related queries (this number is by definition zero for the focal queries.) Whenever we consider a running tally of searches in our estimation below, we consider the period between 2022-01-24 and 2023-01-23 (when we see all searches) and add to the searches that have occurred by any given date during this period the searches that occurred in the prior year, i.e., from 2021-01-24 to 2022-01-23. However, we cannot account for searches even further in the past due to Bing’s retention policy for query data. Finally, we have access to Bing’s internal distance measure between the related queries and the original focal query.

As before, we use the fitted estimates $\hat{r}_{qt}$ using just the $\delta_{q,u(q,t)}$ fixed effect from regression (A9) as our dependent variable in the estimation of (36). Computationally speaking, we obtain our estimates of (36) via a non-linear least squares procedure and standard errors from a block bootstrap (where a block is a focal query.) As our non-linearity correction from Section 6.2 did not yield substantively different estimates there, we avoid implementing a more complicated version of this procedure in this robustness check and simply report the parameters estimated via our demeaned non-linear least squares procedure, noting that these estimates should be interpreted with caution.

As before, we calibrate the intercept α such that the average predicted CTR matches that from our experiment. However, a complication emerges: to take this average, we need to know for each query in the query frequency distribution how many views there are on related queries. While we have this information for the new queries (used in estimation above), we do not have this information for all queries that Bing sees. Hence, we need to predict the value of the term in parentheses, i.e., $\sum_{q' \in Q} f(d(q, q'); \gamma) \times \text{searches}_{q't}$ from just the number of views on the focal query. We use the model

$$\log\left(\sum_{q' \in Q} f(d(q, q'); \gamma) \times \text{searches}_{q't}\right) = \beta_0 + \beta_1 \log(\text{searches}_q) + u_q \quad (38)$$

We fit this equation on our sample of new queries (for which we observe views on related queries), and find $\hat{\beta}_0 = 0.0509(0.0058)$ and $\hat{\beta}_1 = 0.9951(0.0010)$ with $R^2 = 0.94$, suggesting that we can predict this quantity very well. We can thus use

$$r_{qt} = \alpha_q + \beta \frac{1}{1-\theta} (\exp(\beta_0 + \beta_1 \log(n_q)))^{1-\theta} + \varepsilon_{qt}$$

to predict CTR from just an observation of the number of views on a particular (focal) query. This allows us to calibrate α .

We exhibit our results in Table A8. Most importantly, we still find returns that are essentially logarithmic, though once taking into account spillovers, the returns are slightly more convex than logarithmic (i.e., $\hat{\theta} < 1$). Furthermore, we can strongly reject the null hypothesis that additional searches have no impact on performance (i.e., β is significantly different from zero.) Finally, our estimate of γ suggests a limited role of cross-query learning. This is illustrated by Figure A13, which plots the implied weight of searches on related queries against their distance from the focal query. In particular, the horizontal axis measures the distance to the focal query in units of Bing’s internal distance metric; these units are restricted to lie between zero (only assigned for identical queries) and two (practically never assigned.) The solid black line

Table A8: **Cross-Query Economies of Scale Estimates**

Description	Parameter	Estimate	SE
Click-through rate at inception	α	0.1744	-
Value of additional data	β	0.0056	(0.0007)
Shape of returns from data	θ	0.9272	(0.0292)
Relative weight on related queries	γ	3.8327	(0.6659)

Notes: This table provides the estimates of the parameters in (36), obtained via non-linear least squares. Standard errors are from a block-bootstrap clustered at the focal query level.

indicates the weight that our estimates imply for views on a query at a certain distance from a focal query: for instance, at distances of 0.01 our estimates imply a weight of about 0.4, suggesting that each search for a related query at this distance is worth about 40 percent of a search for the original query when it comes to learning how to rank search results. As query distances are hard to interpret, we exhibit the distribution of distances between focal queries and their top-related query (in blue) or their tenth-most related query (in red). We can see that it is rare for queries to have a related query at distance low enough to be assigned a significant weight.

According to the estimates in Table A8, if Bing were to increase its market share by multiplying it by 4.28, its CTR would improve from 23.50 percent to 24.99 percent, an improvement of 1.49 percentage points. As expected, this increase is slightly larger than the 1.29 percentage points that we found for an increase in market share by multiplying it by 4.28 in the main text. In other words, accounting for cross-query spillovers slightly raises our estimates of the importance of economies of scale, but does not lead to any changes in qualitative conclusions.

D Counterfactuals Appendix

Consider the utility of agent i in some counterfactual $\mathcal{C}$. The difference in the user’s perceived utilities—the utility that drives choices—can be written as

$$\Delta u_{i,\mathcal{C}} = \Delta v_{\mathcal{C}} + \Delta b_{\mathcal{C}} + \Delta \chi_i,$$

where $\Delta v_{\mathcal{C}}$ denotes differences in true mean utilities, $\Delta b_{\mathcal{C}}$ denotes additional differences due to misperceptions, and $\Delta \chi_i$ denotes differences in idiosyncratic preferences. To illustrate these terms, we now consider what they look like in the Status Quo. The term for true utilities is $\Delta v_{\mathcal{C}} = \Delta \tilde{\zeta} - \sigma(1 - \delta)$ to account for the difference in the quality of the search engine and for the switching cost. The bias term is $\Delta b_{\mathcal{C}} = \tilde{\zeta}_{-d}^* - \tilde{\zeta}_{-d}$ since the user is not aware of the true quality of the alternative search engine.

To account for inattention, a fraction f_A of users are attentive and make a choice based on $\Delta u_{i,\mathcal{C}}$. In the status quo, $f_A = 1 - \phi$, and in counterfactuals in which all users make an active choice, f_A is simply 1. A fraction f_d of users are inattentive and stay with the default search engine, and a fraction f_{-d} are inattentive and stay with the alternative search engine. Thus, $f_A + f_d + f_{-d} = 1$. In the status quo, for instance, $f_d = \phi$

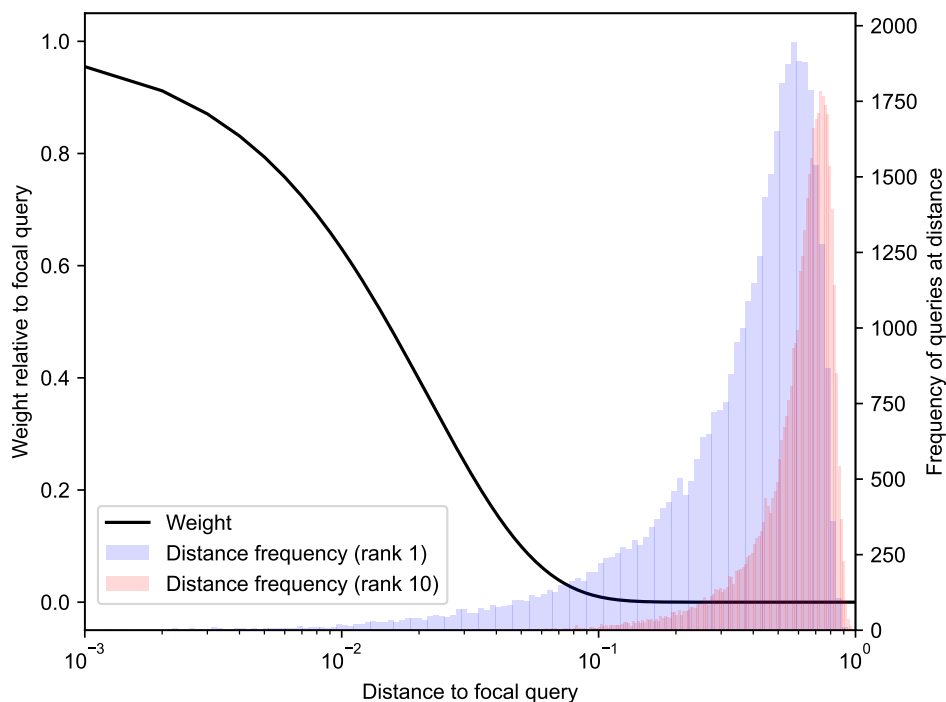


Figure A13: **Cross-Query Learning** Illustrated

Notes: This graph illustrates the (limited) extent of cross-query learning implied by our parameter estimates. The horizontal axis measures the distance to the focal query in units of Bing’s internal distance metric; these units are restricted to lie between zero (only assigned for identical queries) and two (practically never assigned.) The solid black line indicates the weight that our estimates imply for views on a query at a certain distance from a focal query: for instance, at distances of 0.01 our estimates imply a weight of about 0.4, suggesting that each view on a related query at this distance is worth about 40 percent of a view on the original query when it comes to learning how to rank search results. As query distances are hard to interpret, we exhibit the distribution of distances between focal queries and their top-related query (in blue) or their tenth-most related query (in red). We can see that it is rare for queries to have a related query at distance low enough to be assigned a significant weight.

and $f_{-d} = 0$. In counterfactuals in which Bing is the default on all browsers, $f_d = 0$ and $f_{-d} = \phi$ for Chrome users and $f_d = \phi$ and $f_{-d} = 0$ for Edge users.

The market share of search engine $-d$ is given by

$$s_{-d,\mathcal{C}} = f_{-d} + f_A S(\Delta v_{\mathcal{C}} + \Delta b_{\mathcal{C}}),$$

where $S(\cdot)$ is the CDF of the difference in idiosyncratic preferences.

We now derive expressions for consumer surplus. We first derive an expression for attentive consumers. In the absence of any misperceptions, the consumer surplus relative to the utility of the original search engine is given by $\frac{1}{\eta} V(\Delta v_{\mathcal{C}})$, where $V(x) = \int_{-\infty}^x S(x') dx'$. The consumer surplus relative to the utility of the original search engine *in the Status Quo* for attentive consumers is given by

$$\frac{1}{\eta} [V(\Delta v_{\mathcal{C}} + \Delta b_{\mathcal{C}}) - \Delta b_{\mathcal{C}} \cdot S(\Delta v_{\mathcal{C}} + \Delta b_{\mathcal{C}}) + (v_{d,\mathcal{C}} - v_{d,SQ})]$$

The first term in this expression gives the consumer surplus, relative to the utility of the original search engine in $\mathcal{C}$, if users's true utility is indeed described by their perceived utility. The second term is an adjustment due to biases that lead to suboptimal choices. For the share $S(\Delta v_{\mathcal{C}} + \Delta b_{\mathcal{C}})$ of users that choose the alternative search engine, their true utility is lower (or higher) to the extent that they have biases, $-\Delta b_{\mathcal{C}}$. The final term accounts for the fact that we want to measure utility relative to the Status Quo to be able to compare consumer surplus across counterfactuals. We therefore must adjust consumer surplus by the degree to which utility at the anchor point—the original search engine—changed relative to the status quo, $v_{d,\mathcal{C}} - v_{d,SQ}$.

If we now account for permanent inattention, we obtain the following expression for consumer surplus:

$$CS_{\mathcal{C}} = \frac{f_A}{\eta} [V(\Delta v_{\mathcal{C}} + \Delta b_{\mathcal{C}}) - \Delta b_{\mathcal{C}} \cdot S(\Delta v_{\mathcal{C}} + \Delta b_{\mathcal{C}})] + \frac{f_{-d}}{\eta} [\Delta v_{\mathcal{C}} + \mu] + \frac{1}{\eta} (v_{d,\mathcal{C}} - v_{d,SQ})$$

The first term is just the fraction of attentive consumers times their consumer surplus (note that the term $v_{d,\mathcal{C}} - v_{d,SQ}$ affects all users so it is accounted for by the last term in the expression). The second term accounts for the surplus of the fraction f_{-d} of inattentive users that stay with the non-original search engine, whose utility is $\Delta v_{\mathcal{C}}$ —the difference in true utilities—plus the expectation of the error term $\mu = E[\Delta \chi_i]$. For the lognormal error that we use, $\mu = 1 - \exp(\gamma^2/2)$. There is no term for the fraction f_d of users that are stuck with the original search engine because their utility relative to the utility of using the original search engine is simply zero.

D.1 Direct Effects

We now give expressions for $\Delta v_{\mathcal{C}}$, $\Delta b_{\mathcal{C}}$, and $(v_{d,\mathcal{C}} - v_{d,SQ})$ in each of our counterfactuals. Note that the Data Sharing and Data Sharing + Choice Screen counterfactuals are only relevant in equilibrium, since they involve a change in the quality of search engines that arises from the use of data.

Status Quo For the status quo, $\Delta v_{SQ} = \Delta \zeta - \sigma(1 - \delta)$ to account for differences in utilities and for switching costs. The bias term $\Delta b_{SQ} = \tilde{\zeta}_{-d} - \zeta_{-d}^*$ accounts for misperceptions. There are $f_A = 1 - \phi$ attentive users, and $f_d = \phi$ users are inattentive and stay with the default search engine. Trivially, $(v_{d,SQ} - v_{d,SQ}) = 0$.

No Frictions Since there are no switching costs, $\Delta v_{NF} = \Delta \zeta$. Since there are no biases, $\Delta b_{NF} = 0$. There are no inattentive users, so $f_A = 1$. And the true quality of the default search engine is unchanged, so $(v_{d,NF} - v_{d,SQ}) = 0$.

Active Choice Since there are no switching costs, $\Delta v_{AC} = \Delta \zeta$. Customers still have misperceptions about search engines, so $\Delta b_{AC} = \tilde{\zeta}_{-d} - \zeta_{-d}^*$. There are no inattentive users, so $f_A = 1$. And the true quality of the default search engine is unchanged, so $(v_{d,NF} - v_{d,SQ}) = 0$.

Correct Perceptions True utilities do not change, so $\Delta v_{CP} = \Delta \zeta - \sigma(1 - \delta)$. There are no misperceptions, so $\Delta b_{CP} = 0$. There are $f_A = 1 - \phi$ attentive users, and $f_d = \phi$ users are inattentive and stay with the default search engine. And the true quality of the default search engine is unchanged, so $(v_{d,CP} - v_{d,SQ}) = 0$.

Choice Screen Edge users behave as in the Status Quo. Chrome users behave as in the Active Choice counterfactual.

Bing Default Edge users behave just as in the Status Quo. Among Chrome users, switching costs are still present but they go the other way around, so $\Delta v_{BD} = \Delta \zeta + \sigma(1 - \delta)$. We assume Chrome users have no misperceptions about search engines, so $\Delta b_{BD} = 0$. There are $f_A = 1 - \phi$ attentive users, and $f_{-d} = \phi$ users are inattentive and stay with the *alternative* search engine. Finally, the true quality of the default search engine changes because it is now subject to switching costs, so $(v_{d,BD} - v_{d,SQ}) = -\sigma(1 - \delta)$.

Bing Default + Delayed Choice Screen During the first two weeks, the market behaves just as in the Bing Default counterfactual. Starting in week 3, Edge users behave as in the Status Quo. Chrome users behave as in the Correct Perceptions counterfactual. We compute market shares and consumer surplus over a span of six years: we give weight 2/312 to the first two weeks, and we give weight 310/312 to the market after week 3.

Bing Payments Utilities change due to payments, so $\Delta v_{BP} = \eta \Delta p + \Delta \zeta - \sigma(1 - \delta)$ (note that Δp is positive for Chrome users but negative for Edge users). Customers still have the same biases as in the status quo, so $\Delta b_{BP} = \tilde{\zeta}_{-d} - \zeta_{-d}^*$. There are $f_A = 1 - \phi$ attentive users, and $f_d = \phi$ users are inattentive and stay with the default search engine. For Edge users, we must account for the fact that the utility of using Bing increased by \$10, so $v_{d,BP} - v_{d,SQ} = \eta \Delta p$. For Chrome users, the true quality of the default search engine is unchanged, so $v_{d,BP} - v_{d,SQ} = 0$.

D.2 Equilibrium Effects

To account for equilibrium effects, we must account for the fact that true qualities ζ_j are now a function of the share of people using search engines. Suppose that a share $\bar{s}_j$ of people use search engine j across all browsers. Following our economies of scale model, the true qualities are then given by

$$\zeta_j(\bar{s}_j) = \rho \left[\alpha - \frac{\beta}{1-\theta} + \left(\frac{\bar{s}_j}{\bar{s}_{j,SQ}} \right)^{1-\theta} (\hat{r}_j - \alpha + \frac{\beta}{1-\theta}) \right] + \xi_j$$

We can thus derive the following expressions for $\Delta v_{\mathcal{C}}$, $\Delta b_{\mathcal{C}}$, and $(v_{d,\mathcal{C}} - v_{d,SQ})$ in equilibrium (the fractions f_A , f_d , and f_{-d} are the same as in the direct counterfactuals, unless otherwise noted):

No Frictions Since there are no switching costs, $\Delta v_{NF} = \Delta \zeta = \zeta_{-d}^*(\bar{s}_{-d,NF}) - \zeta_d^*(1 - \bar{s}_{-d,NF})$. Since there are no biases, $\Delta b_{NF} = 0$. The true quality of the default search engine changes with the new market shares, so $v_{d,NF} - v_{d,SQ} = \zeta_d^*(1 - \bar{s}_{-d,NF}) - \zeta_d^*(1 - \bar{s}_{-d,SQ})$.

Active Choice Since there are no switching costs, $\Delta v_{AC} = \Delta \zeta = \zeta_{-d}^*(\bar{s}_{-d,AC}) - \zeta_d^*(1 - \bar{s}_{-d,AC})$. Customers still have misperceptions about search engines, so $\Delta b_{AC} = \tilde{\zeta}_{-d} - \zeta_{-d}(\bar{s}_{-d,AC})$. The true quality of the default search engine changes with the new market shares, so $v_{d,AC} - v_{d,SQ} = \zeta_d^*(1 - \bar{s}_{-d,AC}) - \zeta_d^*(1 - \bar{s}_{-d,SQ})$.

Correct Perceptions True utilities change due to new market shares, so $\Delta v_{CP} = \zeta_{-d}^*(\bar{s}_{-d,CP}) - \zeta_d^*(1 - \bar{s}_{-d,CP}) - \sigma(1 - \delta)$. There are no misperceptions, so $\Delta b_{CP} = 0$. The true quality of the default search engine changes with the new market shares, so $v_{d,CP} - v_{d,SQ} = \zeta_d^*(1 - \bar{s}_{-d,CP}) - \zeta_d^*(1 - \bar{s}_{-d,SQ})$.

No Frictions + Data Sharing All expressions are the same as in No Frictions except that the quality of Bing is given by $\zeta_{-d}^*(1)$ instead of $\zeta_{-d}^*(\bar{s}_{-d,NF})$.

Choice Screen Edge users behave as in the Status Quo, and Chrome users behave as in the Active Choice counterfactual, following the expressions above. Qualities must be adjusted to account for equilibrium effects: they are now $\zeta_{-d}^*(\bar{s}_{-d,CS})$ and $\zeta_d^*(1 - \bar{s}_{-d,CS})$.

Bing Default Edge users behave just as in the Status Quo, but qualities do change because of economies of scale: $\Delta v_{BD,E} = \Delta \tilde{\zeta} - \sigma(1 - \delta)$, $\Delta b_{BD,E} = \tilde{\zeta}_{-d} - \zeta_{-d}^*(\bar{s}_{-d,BD})$, and $(v_{d,SQ} - v_{d,SQ}) = 0$. The true quality of Bing changes: $(v_{d,BD,E} - v_{d,BD,E}) = \zeta_d^*(1 - \bar{s}_{-d,BD}) - \zeta_d^*(1 - \bar{s}_{-d,SQ})$. For Chrome users, switching costs are still present but they go the other way around: $\Delta v_{BD,C} = \zeta_{-d}^*(\bar{s}_{-d,BD}) - \zeta_d^*(1 - \bar{s}_{-d,BD}) + \sigma(1 - \delta)$. Chrome users have no misperceptions so $\Delta b_{BD,C} = 0$, and the true quality of the default search engine changes due to the change in market shares and because it is now subject to switching costs so $(v_{d,BD,C} - v_{d,BD,C}) = \zeta_d^*(1 - \bar{s}_{-d,BD}) - \zeta_d^*(1 - \bar{s}_{-d,SQ}) - \sigma(1 - \delta)$.

Bing Default + Delayed Choice Screen During the first two weeks, the market behaves just as in the Bing Default counterfactual. Starting on week 3, Edge users behave as in the Status Quo. Chrome users behave as in the Correct Perceptions counterfactual: $\Delta v_{DCS,C} = \zeta_{-d}^*(\bar{s}_{-d,DCS}) - \zeta_d^*(1 - \bar{s}_{-d,DCS}) - \sigma(1 - \delta)$, $\Delta b_{DCS,C} = 0$, and $(v_{d,DCS,C} - v_{d,SQ}) = \zeta_d^*(1 - \bar{s}_{-d,DCS}) - \zeta_d^*(1 - \bar{s}_{-d,SQ})$. We compute market shares and consumer surplus using the same weights as in the direct effects counterfactual.

Bing Payments Utilities change due to payments. Thus, $\Delta v_{BP} = \eta \Delta p + \zeta_{-d}^*(\bar{s}_{-d,BP}) - \zeta_d^*(1 - \bar{s}_{-d,BP}) - \sigma(1 - \delta)$. Customers still have the same biases as in the status quo, so $\Delta b_{BP} = \tilde{\zeta}_{-d} - \zeta_{-d}^*(\bar{s}_{-d,BP})$. The true quality of the default search engine changes with the new market shares, so $v_{d,BP} - v_{d,SQ} = \zeta_d^*(1 - \bar{s}_{-d,BP}) - \zeta_d^*(1 - \bar{s}_{-d,SQ})$ for Chrome users. For Edge users, we also need to account for the fact that the utility of using Bing changed by \$10, so $v_{d,BP} - v_{d,SQ} = \zeta_d^*(1 - \bar{s}_{-d,BP}) - \zeta_d^*(1 - \bar{s}_{-d,SQ}) + \eta \Delta p$.

Data Sharing True qualities are $\zeta_G^*(\bar{s}_G)$ and $\zeta_B^*(1)$. True utilities are $\Delta v_{DS} = \zeta_{-d}^* - \zeta_d^* - \sigma(1 - \delta)$. The bias term $\Delta b_{DS} = \tilde{\zeta}_{-d} - \zeta_{-d}^*$ is as in the Status Quo. The true quality of the default search engine changes because of data sharing, so $v_{G,DS} - v_{G,SQ} = \zeta_G^*(\bar{s}_{G,DS}) - \zeta_G^*(\bar{s}_{G,SQ})$ and $v_{B,DS} - v_{B,SQ} = \zeta_B^*(1) - \zeta_B^*(\bar{s}_{B,SQ})$.

D.2.1 Computing equilibria

In each of the above counterfactuals, we can plug in the above expressions into our expression for market shares to obtain the following expression for market shares among users of browser b :

$$s_{b,-d} = S(\Delta v_{b,\mathcal{C}}(\bar{s}_{-d}) + \Delta b_{b,\mathcal{C}}(\bar{s}_{-d})).$$

We can aggregate those market shares to obtain Google's total market share

$$\bar{s}_G = \frac{n_{CH}S(\Delta v_{CH,\mathcal{C}}(1 - \bar{s}_G) + \Delta b_{CH,\mathcal{C}}(1 - \bar{s}_G)) + n_{ED}S(\Delta v_{ED,\mathcal{C}}(\bar{s}_G) + \Delta b_{ED,\mathcal{C}}(\bar{s}_G))}{n_{CH} + n_{ED}},$$

where n_{CH} and n_{ED} represents the number of users on Chrome and Edge, respectively.

Finding an equilibrium consists of computing a solution $\bar{s}_{G,\mathcal{C}}$ to the above equation. We implement this using the bisection method with shares one and zero as starting points. Once we obtain such solution, it is straightforward to compute qualities, which we can use to compute $\Delta v_{\mathcal{C}}$, $\Delta b_{\mathcal{C}}$, and $(v_{d,\mathcal{C}} - v_{d,SQ})$ as well as equilibrium market shares and consumer surplus.

D.3 Additional Counterfactual Simulation Results

Table A9: Counterfactual Simulations: Strong Effect on Search Result Relevance

Description	Combined		Chrome		Edge	
	(1) Google share (%)	(2) CS gain (\$/user-year)	(3) Google share (%)	(4) CS gain (\$/user-year)	(5) Google share (%)	(6) CS gain (\$/user-year)
<i>Benchmarks</i>						
Status Quo	88.9	0.00	98.7	0.00	22.4	0.00
No Frictions	75.3	5.93	81.5	0.60	33.8	41.87
Active Choice	89.1	5.47	97.3	0.08	33.8	41.73
Correct Perceptions	79.5	0.32	87.9	0.35	22.4	0.14
No Frictions + Data Sharing	74.8	6.05	80.8	0.68	33.8	42.18
<i>Policy Interventions</i>						
Choice Screen	87.6	0.07	97.3	0.08	22.4	0.03
Bing Default	49.5	-72.80	53.5	-83.65	22.4	0.32
Bing Default + Delayed Choice Screen	73.8	0.07	81.5	0.05	22.4	0.19
Bing Payments (\$10)	51.7	107.48	56.8	92.30	17.2	209.71
Data Sharing	88.9	0.08	98.7	0.01	22.4	0.56

Notes: This table presents the equilibrium effects of the counterfactual simulation results from the procedure described in Section 7 under the scenario where all consumers experience a strong effect on search result relevance. Results are computed exactly as in Panel B of Table 7, except that, rather than using our point estimate for quality response, we use the lower bound of the 95% confidence interval.