

The search engine manipulation effect (SEME) and its possible impact on the outcomes of elections

Robert Epstein¹ and Ronald E. Robertson

American Institute for Behavioral Research and Technology, Vista, CA 92084

Edited by Jacob N. Shapiro, Princeton University, Princeton, NJ, and accepted by the Editorial Board July 8, 2015 (received for review October 16, 2014)

Internet search rankings have a significant impact on consumer choices, mainly because users trust and choose higher-ranked results more than lower-ranked results. Given the apparent power of search rankings, we asked whether they could be manipulated to alter the preferences of undecided voters in democratic elections. Here we report the results of five relevant double-blind, randomized controlled experiments, using a total of 4,556 undecided voters representing diverse demographic characteristics of the voting populations of the United States and India. The fifth experiment is especially notable in that it was conducted with eligible voters throughout India in the midst of India's 2014 Lok Sabha elections just before the final votes were cast. The results of these experiments demonstrate that (i) biased search rankings can shift the voting preferences of undecided voters by 20% or more, (ii) the shift can be much higher in some demographic groups, and (iii) search ranking bias can be masked so that people show no awareness of the manipulation. We call this type of influence, which might be applicable to a variety of attitudes and beliefs, the search engine manipulation effect. Given that many elections are won by small margins, our results suggest that a search engine company has the power to influence the results of a substantial number of elections with impunity. The impact of such manipulations would be especially large in countries dominated by a single search engine company.

search engine manipulation effect | search rankings | Internet influence | voter manipulation | digital bandwagon effect

Recent research has demonstrated that the rankings of search results provided by search engine companies have a dramatic impact on consumer attitudes, preferences, and behavior (1–12); this is presumably why North American companies now spend more than 20 billion US dollars annually on efforts to place results at the top of rankings (13, 14). Studies using eye-tracking technology have shown that people generally scan search engine results in the order in which the results appear and then fixate on the results that rank highest, even when lower-ranked results are more relevant to their search (1–5). Higher-ranked links also draw more clicks, and consequently people spend more time on Web pages associated with higher-ranked search results (1–9). A recent analysis of ~300 million clicks on one search engine found that 91.5% of those clicks were on the first page of search results, with 32.5% on the first result and 17.6% on the second (7). The study also reported that the bottom item on the first page of results drew 140% more clicks than the first item on the second page (7). These phenomena occur apparently because people trust search engine companies to assign higher ranks to the results best suited to their needs (1–4, 11), even though users generally have no idea how results get ranked (15).

Why do search rankings elicit such consistent browsing behavior? Part of the answer lies in the basic design of a search engine results page: the list. For more than a century, research has shown that an item's position on a list has a powerful and persuasive impact on subjects' recollection and evaluation of that item (16–18). Specific order effects, such as primacy and recency, show that the first and last items presented on a list, respectively, are more likely to be recalled than items in the middle (16, 17).

Primacy effects in particular have been shown to have a favorable influence on the formation of attitudes and beliefs (18–20), enhance perceptions of corporate performance (21), improve ratings of items on a survey (22–24), and increase purchasing behavior (25). More troubling, however, is the finding that primacy effects have a significant impact on voting behavior, resulting in more votes for the candidate whose name is listed first on a ballot (26–32). In one recent experimental study, primacy accounted for a 15% gain in votes for the candidate listed first (30). Although primacy effects have been shown to extend to hyperlink clicking behavior in online environments (33–35), no study that we are aware of has yet examined whether the deliberate manipulation of search engine rankings can be leveraged as a form of persuasive technology in elections. Given the power of order effects and the impact that search rankings have on consumer attitudes and behavior, we asked whether the deliberate manipulation of search rankings pertinent to candidates in political elections could alter the attitudes, beliefs, and behavior of undecided voters.

It is already well established that biased media sources such as newspapers (36–38), political polls (39), and television (40) sway voters (41, 42). A 2007 study by DellaVigna and Kaplan found, for example, that whenever the conservative-leaning Fox television network moved into a new market in the United States, conservative votes increased, a phenomenon they labeled the Fox News Effect (40). These researchers estimated that biased coverage by Fox News was sufficient to shift 10,757 votes in Florida during the 2000 US Presidential election: more than enough to flip the deciding state in the election, which was carried by the Republican presidential candidate by only 537 votes. The Fox News Effect was also found to be smaller in television markets that were more competitive.

We believe, however, that the impact of biased search rankings on voter preferences is potentially much greater than the influence of traditional media sources (43), where parties compete in

Significance

We present evidence from five experiments in two countries suggesting the power and robustness of the search engine manipulation effect (SEME). Specifically, we show that (i) biased search rankings can shift the voting preferences of undecided voters by 20% or more, (ii) the shift can be much higher in some demographic groups, and (iii) such rankings can be masked so that people show no awareness of the manipulation. Knowing the proportion of undecided voters in a population who have Internet access, along with the proportion of those voters who can be influenced using SEME, allows one to calculate the win margin below which SEME might be able to determine an election outcome.

Author contributions: R.E. and R.E.R. designed research, performed research, contributed new reagents/analytic tools, analyzed data, and wrote the paper.

The authors declare no conflict of interest.

This article is a PNAS Direct Submission. J.N.S. is a guest editor invited by the Editorial Board.

Freely available online through the PNAS open access option.

¹To whom correspondence should be addressed. Email: re@aibr.org.

This article contains supporting information online at www.pnas.org/lookup/suppl/doi:10.1073/pnas.1419828112/-DCSupplemental.

an open marketplace for voter allegiance. Search rankings are controlled in most countries today by a single company. If, with or without intervention by company employees, the algorithm that ranked election-related information favored one candidate over another, competing candidates would have no way of compensating for the bias. It would be as if Fox News were the only television network in the country. Biased search rankings would, in effect, be an entirely new type of social influence, and it would be occurring on an unprecedented scale. Massive experiments conducted recently by social media giant Facebook have already introduced other unprecedented types of influence made possible by the Internet. Notably, an experiment reported recently suggested that flashing “VOTE” ads to 61 million Facebook users caused more than 340,000 people to vote that day who otherwise would not have done so (44). Zittrain has pointed out that if Facebook executives chose to prompt only those people who favored a particular candidate or party, they could easily flip an election in favor of that candidate, performing a kind of “digital gerrymandering” (45).

We evaluated the potential impact of biased search rankings on voter preferences in a series of experiments with the same general design. Subjects were asked for their opinions and voting preferences both before and after they were allowed to conduct research on candidates using a mock search engine we had created for this purpose. Subjects were randomly assigned to groups in which the search results they were shown were biased in favor of one candidate or another, or, in a control condition, in favor of neither candidate. Would biased search results change the opinions and voting preferences of undecided voters, and, if so, by how much? Would some demographic groups be more vulnerable to such a manipulation? Would people be aware that they were viewing biased rankings? Finally, what impact would familiarity with the candidates have on the manipulation?

Study 1: Three Experiments in San Diego, CA

To determine the potential for voter manipulation using biased search rankings, we initially conducted three laboratory-based experiments in the United States, each using a double-blind control group design with random assignment. For each of the experiments, we recruited 102 eligible voters through newspaper and online advertisements, as well through notices in senior recreation centers, in the San Diego, CA, area.* The advertisements offered USD\$25 for each subject's participation, and subjects were prescreened in an attempt to match diverse demographic characteristics of the US voting population (46).

Each of the three experiments used 30 actual search results and corresponding Web pages relating to the 2010 election to determine the prime minister of Australia. The candidates were Tony Abbott and Julia Gillard, and the order in which their names were presented was counterbalanced in all conditions. This election was used to minimize possible preexisting biases by US study participants and thus to try to guarantee that our subjects would be truly “undecided.” In each experiment, subjects were randomly assigned to one of three groups: (i) rankings favoring Gillard (which means that higher-ranked search results linked to Web pages that portrayed Gillard as the better candidate), (ii) rankings favoring Abbott, or (iii) rankings favoring neither (Fig. 1 A–C). The order of these rankings was determined based on ratings of Web pages provided by three independent observers. Neither the subjects nor the research assistants who supervised them knew either the hypothesis of the experiment or the groups to which subjects were assigned.

Initially, subjects read brief biographies of the candidates and rated them on 10-point Likert scales with respect to their overall impression of each candidate, how much they trusted each candidate, and how much they liked each candidate. They were

also asked how likely they would be to vote for one candidate or the other on an 11-point scale ranging from –5 to +5, as well as to indicate which of the two candidates they would vote for if the election were held that day.

The subjects then spent up to 15 min gathering more information about the candidates using a mock search engine we had created (called Kadoodle), which gave subjects access to five pages of search results with six results per page. As is usual with search engines, subjects could click on any search result to view the corresponding Web page, or they could click on numbers at the bottom of each results page to view other results pages. The same search results and Web pages were used for all subjects in each experiment; only the order of the search results was varied (Fig. 1). Subjects had the option to end the search whenever they felt they had acquired sufficient information to make a sound decision. At the conclusion of the search, subjects rated the candidates again. When their ratings were complete, subjects were asked (on their computer screens) whether anything about the search rankings they had viewed “bothered” them; they were then given an opportunity to write at length about what, if anything, had bothered them. We did not ask specifically whether the search rankings appeared to be “biased” to avoid false positives typically generated by leading or suggestive questions (47).

Regarding the ethics of our study, our manipulation could have no impact on a past election, and we were also not concerned that it could affect the outcome of future elections, because the number of subjects we recruited was small and, to our knowledge, included no Australian voters. Moreover, our study was designed so that it did not favor any one candidate, so there was no overall bias. The study presented no more than minimal risk to subjects and was approved by the Institutional Review Board (IRB) of the American Institute for Behavioral Research and Technology (AIBRT). Informed consent was obtained from all subjects.

In aggregate for the first three experiments in San Diego, CA, the demographic characteristics of our subjects (mean age, 42.5 y; SD = 18.1 y; range, 18–95 y) did not differ from characteristics of the US voting population by more than the following

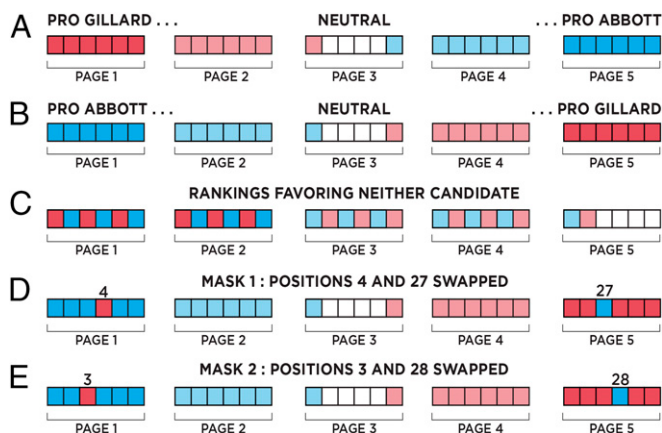


Fig. 1. Search rankings for the three experiments in study 1. (A) For subjects in group 1 of experiment 1, 30 search results that linked to 30 corresponding Web pages were ranked in a fixed order that favored candidate Julia Gillard, as follows: those favoring Gillard (from highest to lowest rated pages), then those favoring neither candidate, then those favoring Abbott (from lowest to highest rated pages). (B) For subjects in group 2 of experiment 1, the search results were displayed in precisely the opposite order so that they favored the opposing candidate, Tony Abbott. (C) For subjects in group 3 of experiment 1 (the control group), the ranking favored neither candidate. (D) For subjects in groups 1 and 2 of experiment 2, the rankings bias was masked slightly by swapping results that had originally appeared in positions 4 and 27. Thus, on the first page of search results, five of the six results—all but the one in the fourth position—favored one candidate. (E) For subjects in groups 1 and 2 of experiment 3, a more aggressive mask was used by swapping results that had originally appeared in positions 3 and 28.

*Although all participants claimed to be eligible voters in the prescreening, we later discovered that 6.9% of subjects marked “I don’t know” and 5.2% of subjects marked “No” in response to a question asking “If you are not currently registered, are you eligible to register for elections?”

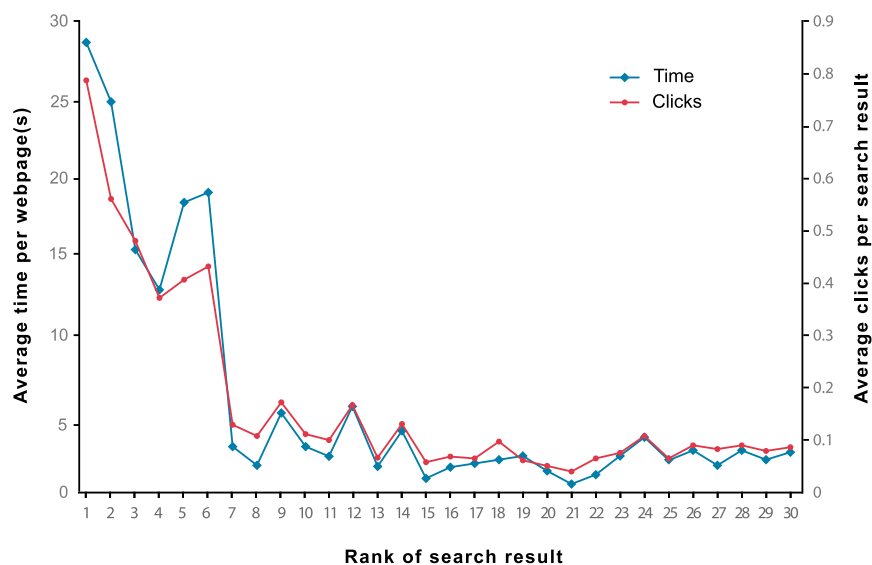


Fig. 2. Clicks on search results and time allocated to Web pages as a function of search result rank, aggregated across the three experiments in study 1. Subjects spent less time on Web pages corresponding to lower-ranked search results (blue curve) and were less likely to click on lower-ranked results (red curve). This pattern is found routinely in studies of Internet search engine use (1–12).

margins: 6.4% within any category of the age or sex measures; 14.1% within any category of the race measure; 18.7% within any category of the income or education measures; and 21.1% within any category of the employment status measure (Table S1). Subjects' political inclinations were fairly balanced, with 20.3% identifying themselves as conservative, 28.8% as moderate, 22.5% as liberal, and 28.4% as indifferent. Political party affiliation, however, was less balanced, with 21.6% identifying as Republican, 19.6% as Independent, 44.8% as Democrat, 6.2% as Libertarian, and 7.8% as other. In aggregate, subjects reported conducting an average of 7.9 searches (SD = 17.5) per day using search engines, and 52.3% reported having conducted searches to learn about political candidates. They also reported having little or no familiarity with the candidates (mean familiarity on a scale of 1–10, 1.4; SD = 0.99). On average, subjects in the first three experiments spent 635.9 s (SD = 307.0) using our mock search engine.

As expected, higher search rankings drew more clicks, and the pattern of clicks for the first three experiments correlated strongly with the pattern found in a recent analysis of ~300 million clicks [$r(13) = 0.90$, $P < 0.001$; Kolmogorov–Smirnov test of differences in distributions: $D = 0.033$, $P = 0.31$; Fig. 2] (7). In addition, subjects spent more time on Web pages associated with higher-ranked results (Fig. 2), as well as substantially more time on earlier search pages (Fig. 3).

In experiment 1, we found no significant differences among the three groups with respect to subjects' ratings of the candidates before Web research (Table S2). Following the Web research, all candidate ratings in the bias groups shifted in the predicted directions compared with candidate ratings in the control group (Table 1).

Before Web research, we found no significant differences among the three groups with respect to the proportions of people who said that they would vote for one candidate or the other if the election were held today (Table 2). Following Web research, significant differences emerged among the three groups for this measure (Table 2), and the number of subjects who said they would vote for the favored candidate in the two bias groups combined increased by 48.4% (95% CI, 30.8–66.0%; McNemar's test, $P < 0.01$).

We define the latter percentage as vote manipulation power (VMP). Thus, before the Web search, if a total of x subjects in the bias groups said they would vote for the target candidate, and if, following the Web search, a total of x' subjects in the bias groups said they would vote for the target candidate, $VMP = (x' - x)/x$. The VMP is, we believe, the key measure that an administrator would want to know if he or she were trying to manipulate an election using SEME.

Using a more sensitive measure than forced binary choice, we also asked subjects to estimate the likelihood, on an 11-point

scale from -5 to $+5$, that they would vote for one candidate or the other if the election were held today. Before Web research, we found no significant differences among the three groups with respect to the likelihood of voting for one candidate or the other [Kruskal–Wallis (K–W) test: $\chi^2(2) = 1.384$, $P = 0.501$]. Following Web research, the likelihood of voting for either candidate in the bias groups diverged from their initial scale values by 3.71 points in the predicted directions [Mann–Whitney (M–W) test: $u = 300.5$, $P < 0.01$]. Notably, 75% of subjects in the bias groups showed no awareness of the manipulation. We counted subjects as showing awareness of the manipulation if (i) they had clicked on the box indicating that something bothered them about the rankings and (ii) we found specific terms or phrases in their open-ended comments suggesting that they were aware of bias in the rankings (SI Text).

In experiment 2, we sought to determine whether the proportion of subjects who were unaware of the manipulation could be increased with voter preferences still shifting in the predicted directions. We accomplished this by masking our manipulation to some extent. Specifically, the search result that had appeared in the fourth position on the first page of the search results favoring Abbott in experiment 1 was swapped with the corresponding search result favoring Gillard (Fig. 1D). Before Web research, we found no significant differences among the three

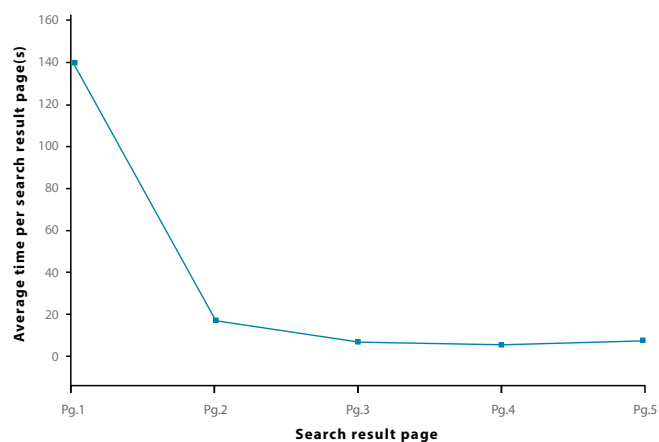


Fig. 3. Amount of time, aggregated across the three experiments in study 1, that subjects spent on each of the five search pages. Subjects spent most of their time on the first search page, a common finding in Internet search engine research (1–12).

Table 1. Postsearch shifts in voting preferences for study 1

Experiment	Candidate	Rating	Mean deviation from control (SE)			
			Gillard bias	<i>u</i>	Abbott bias	<i>u</i>
1	Gillard	Impression	1.44 (0.56)*	761.0	−1.52 (0.56)**	380.5
		Trust	1.26 (0.53)**	779.0	−1.85 (0.48)**	330.5
		Like	0.26 (0.54)	615.5	−1.73 (0.65)**	387.0
	Abbott	Impression	−2.29 (0.73)**	373.0	1.11 (0.72)**	766.5
		Trust	−2.02 (0.63)**	384.0	0.67 (0.76)	679.0
		Like	−1.55 (0.71)	460.5	1.17 (0.64)*	733.0
2	Gillard	Impression	0.97 (0.65)	704.0	−2.38 (0.79)***	325.0
		Trust	0.94 (0.72)	691.5	−2.17 (0.74)**	332.5
		Like	0.55 (0.76)	639.5	−1.82 (0.66)**	378.0
	Abbott	Impression	−1.44 (0.81)*	395.5	1.17 (0.75)*	742.0
		Trust	−0.79 (0.81)	453.5	1.85 (0.72)**	774.5
		Like	−1.44 (0.70)*	429.0	0.64 (0.71)	690.0
3	Gillard	Impression	1.44 (0.73)*	717.5	−0.55 (0.69)	507.5
		Trust	0.47 (0.70)	620.0	−0.23 (0.56)	466.5
		Like	0.44 (0.65)	623.5	−0.41 (0.70)	528.5
	Abbott	Impression	−0.32 (0.70)	534.0	1.26 (0.60)*	750.5
		Trust	−0.73 (0.65)	498.5	1.50 (0.58)**	795.0
		Like	−0.50 (0.61)	496.0	0.88 (0.62)	681.5

* $P < 0.05$, ** $P < 0.01$, and *** $P < 0.001$: Mann–Whitney *u* tests were conducted between the control group and each of the bias groups.

groups with respect to subjects' ratings of the candidates (Table S2). Following the Web research, all candidate ratings in the bias groups shifted in the predicted directions compared with candidate ratings in the control group (Table 1).

Before Web research, we found no significant differences among the three groups with respect to voting proportions (Table 2). Following Web research, significant differences emerged among the three groups for this measure (Table 2), and the VMP was 63.3% (95% CI, 46.1–80.6%; McNemar's test, $P < 0.001$).

For the more sensitive measure (the 11-point scale), we found no significant differences among the three groups with respect to the likelihood of voting for one candidate or the other before Web research [K–W test: $\chi^2(2) = 0.888$, $P = 0.642$]. Following Web research, the likelihood of voting for either candidate in the bias groups diverged from their initial scale values by 4.44 points in the predicted directions (M–W test: $u = 237.5$, $P < 0.001$). In addition, the proportion of people who showed no awareness of the manipulation increased from 75% in experiment 1 to 85% in

experiment 2, although the difference between these percentages was not significant ($\chi^2 = 2.264$, $P = 0.07$).

In experiment 3, we sought to further increase the proportion of subjects who were unaware of the manipulation by using a more aggressive mask. Specifically, the search result that had appeared in the third position on the first page of the search results favoring Abbott in experiment 1 was swapped with the corresponding search result favoring Gillard (Fig. 1E). This mask is a more aggressive one because higher ranked results are viewed more and taken more seriously by people conducting searches (1–12).

Before Web research, we found no significant differences among the three groups with respect to subjects' ratings of candidates (Table S2). Following the Web research, all candidate ratings in the bias groups shifted in the predicted directions compared with candidate ratings in the control group (Table 1).

Before Web research, we found no significant differences among the three groups with respect to voting proportions (Table 2). Following Web research, significant differences did not emerge among

Table 2. Comparison of voting proportions before and after Web research by group for studies 1 and 2

Study	Experiment	Group	Simulated vote before Web research		χ^2	Simulated vote after Web research		χ^2	VMP
			Gillard	Abbott		Gillard	Abbott		
1	1	1	8	26	5.409	22	12	8.870*	48.4%**
		2	11	23		10	24		
		3	17	17		14	20		
	2	1	16	18	2.197	27	7	14.274***	63.3%***
		2	20	14		12	22		
		3	14	20		22	12		
	3	1	17	17	2.199	22	12	3.845	36.7%*
		2	21	13		15	19		
		3	15	19		15	19		
2	4	1	317	383	1.047	489	211	196.280***	37.1%***
		2	316	384		228	472		
		3	333	367		377	323		

McNemar's test was conducted to assess VMP significance. VMP, percent increase in subjects in the bias groups combined who said that they would vote for the favored candidate.

* $P < 0.05$; ** $P < 0.01$; and *** $P < 0.001$: Pearson χ^2 tests were conducted among all three groups.

the three groups for this measure (Table 2); the VMP, however, was 36.7% (95% CI, 19.4–53.9%; McNemar's test, $P < 0.05$).

For the more sensitive measure (the 11-point scale), we found no significant differences among the three groups with respect to the likelihood of voting for one candidate or the other before Web research [K-W test: $\chi^2(2) = 0.624$, $P = 0.732$]. Following Web research, the likelihood of voting for either candidate in the bias groups diverged from their initial scale values by 2.62 points in the predicted directions (M-W test: $u = 297.0$, $P < 0.001$). Notably, in experiment 3, no subjects showed awareness of the rankings bias, and the difference between the proportions of subjects who appeared to be unaware of the manipulations in experiments 1 and 3 was significant ($\chi^2 = 19.429$, $P < 0.001$).

Although the findings from these first three experiments were robust, the use of small samples from one US city limited their generalizability and might even have exaggerated the effect size (48).

Study 2: Large-Scale National Online Replication of Experiment 3

To better assess the generalizability of SEME to the US population at large, we used a diverse national sample of 2,100 individuals[†] from all 50 US states (Table S1), recruited using Amazon's Mechanical Turk (mturk.com), an online subject pool that is now commonly used by behavioral researchers (49, 50). Subjects (mean age, 33.9 y; SD = 11.9 y; range, 18–81 y) were exposed to the same aggressive masking procedure we used in experiment 3 (Fig. 1E). Each subject was paid USD\$1 for his or her participation.

Regarding ethical concerns, as in study 1, our manipulation could have no impact on a past election, and we were not concerned that it could affect the outcome of future elections. Moreover, our study was designed so that it did not favor any one candidate, so there was no overall bias. The study presented no more than minimal risk to subjects and was approved by AIBRT's IRB. Informed consent was obtained from all subjects.

Subjects' political inclinations were less balanced than those in study 1, with 19.5% of subjects identifying themselves as conservative, 24.2% as moderate, 50.2% as liberal, and 6.3% as indifferent; 16.1% of subjects identified themselves as Republican, 29.9% as Independent, 43.2% as Democrat, 8.0% as Libertarian, and 2.9% as other. Subjects reported having little or no familiarity with the candidates (mean, 1.9; SD = 1.7). As one might expect in a study using only Internet-based subjects, self-reported search engine use was higher in study 2 than in study 1 [mean searches per day, 15.3; SD = 26.3; $t(529.5)^{\ddagger} = 6.9$, $P < 0.001$], and more subjects reported having previously used a search engine to learn about political candidates (86.0%, $\chi^2 = 204.1$, $P < 0.001$). Subjects in study 2 also spent less time using our mock search engine [mean total time, 309.2 s; SD = 278.7; $t(381.9)^{\ddagger} = -17.6$, $P < 0.001$], but patterns of search result clicks and time spent on Web pages were similar to those we found in study 1 [clicks: $r(28) = 0.98$, $P < 0.001$; Web page time: $r(28) = 0.98$, $P < 0.001$] and to those routinely found in other studies (1–12).

Before Web research, we found no significant differences among the three groups with respect to subjects' ratings of the candidates (Table S3). Following the Web research, all candidate ratings in the bias groups shifted in the predicted directions compared with candidate ratings in the control group (Table 3).

Before Web research, we found no significant differences among the three groups with respect to voting proportions (Table 2). Following Web research, significant differences emerged among the three groups for this measure (Table 2), and the VMP was 37.1% (95% CI, 33.5–40.7%; McNemar's test, $P < 0.001$). Using poststratification and weights obtained from the 2010 US Census (46) and a 2011 study from Gallup (51), which were scaled to size for age, sex, race, and education, the VMP was 36.7% (95% CI, 33.2–

40.3%; McNemar's test, $P < 0.001$). When weighted using the same demographics via classical regression poststratification (52) (Table S4), the VMP was 33.5% (95% CI, 30.1–37.0%, McNemar's test, $P < 0.001$).

For the more sensitive measure (the 11-point scale), we found no significant differences among the three groups with respect to the likelihood of voting for one candidate or the other before Web research [K-W test: $\chi^2(2) = 2.790$, $P = 0.248$]. Following Web research, the likelihood of voting for either candidate in the bias groups diverged from their initial scale values by 3.03 points in the predicted directions (M-W test: $u = 1.29 \times 10^3$, $P < 0.001$). As one might expect of a more Internet-fluent sample, the proportion of subjects showing no awareness of the manipulation dropped to 91.4%.

The number of subjects in study 1 was too small to look at demographic differences. In study 2, we found substantial differences in how vulnerable different demographic groups were to SEME. Consistent with previous findings on the moderators of order effects (30–32), for example, we found that subjects reporting a low familiarity with the candidates (familiarity less than 5 on a scale from 1 to 10) were more vulnerable to SEME (VMP = 38.7%; 95% CI, 34.9–42.4%; McNemar's test, $P < 0.001$) than were subjects who reported high familiarity with the candidates (VMP = 19.3%; 95% CI, 9.1–29.5%; McNemar's test, $P < 0.05$), and this difference was significant ($\chi^2 = 8.417$, $P < 0.01$).

We found substantial differences in vulnerability to SEME among a number of different demographic groups (SI Text). Although the groups we examined were overlapping and somewhat arbitrary, if one were manipulating an election, information about such differences would have enormous practical value. For example, we found that self-labeled Republicans were more vulnerable to SEME (VMP = 54.4%; 95% CI, 45.2–63.5%; McNemar's test, $P < 0.001$) than were self-labeled Democrats (VMP = 37.7%; 95% CI, 32.3–43.1%; McNemar's test, $P < 0.001$) and that self-labeled divorcees were more vulnerable (VMP = 46.7%; 95% CI, 32.1–61.2%; McNemar's test, $P < 0.001$) than were self-labeled married subjects (VMP = 32.4%; 95% CI, 26.8–38.1%; McNemar's test, $P < 0.001$). Among the most vulnerable groups we identified were Moderate Republicans (VMP = 80.0%; 95% CI, 62.5–97.5%; McNemar's test, $P < 0.001$), whereas among the least vulnerable groups were people who reported a household income of \$40,000 to \$49,999 (VMP = 22.5%; 95% CI, 13.8–31.1%; McNemar's test, $P < 0.001$).

Notably, awareness of the manipulation not only did not nullify the effect, it seemed to enhance it, perhaps because people trust search order so much that awareness of the bias serves to confirm the superiority of the favored candidate. The VMP for people who showed no awareness of the biased search rankings ($n = 1,280$) was 36.3% (95% CI, 32.6–40.1%; McNemar's test, $P < 0.001$), whereas the VMP for people who showed awareness of the bias ($n = 120$) was 45.0% (95% CI, 32.4–57.6%; McNemar's test, $P < 0.001$).

Having now replicated the effect with a large and diverse sample of US subjects, we were concerned about the weaknesses associated with testing subjects on a somewhat abstract election (the election in Australia) that had taken place years before and in which subjects were unfamiliar with the candidates. In real elections, people are familiar with the candidates and are influenced, sometimes on a daily basis, by aggressive campaigning. Presumably, either of these two factors—familiarity and outside influence—could potentially minimize or negate the influence of biased search rankings on voter preferences. We therefore asked if SEME could be replicated with a large and diverse sample of real voters in the midst of a real election campaign.

Study 3: SEME Evaluated During the 2014 Lok Sabha Elections in India

In our fifth experiment, we sought to manipulate the voting preferences of undecided eligible voters in India during the 2014 national Lok Sabha elections there. This election was the largest democratic election in history, with more than 800 million eligible voters and more than 430 million votes ultimately cast. We accomplished this by randomly assigning undecided English-speaking

[†]As in study 1, although all participants claimed to be eligible voters in the prescreening, we later discovered that 4.7% of subjects marked "I don't know" and 2.6% of subjects marked "No" in response to a question asking "If you are not currently registered, are you eligible to register for elections?"

[‡]Degrees of freedom adjusted for significant inequality of variances (Welch's t test).

Table 3. Postsearch shifts in voting preferences for study 2

Candidate	Rating	Mean deviation from control (SE)			
		Gillard bias	<i>u</i>	Abbott bias	<i>u</i>
Gillard	Impression	0.65 (0.10)***	288,299.5	−1.25 (0.12)***	168,203.5
	Trust	0.61 (0.10)***	283,491.0	−1.21 (0.11)***	167,658.5
	Like	0.50 (0.10)***	279,967.0	−1.25 (0.11)***	166,544.0
Abbott	Impression	−0.96 (0.13)***	189,290.5	1.35 (0.12)***	326,067.0
	Trust	−1.09 (0.14)***	183,993.0	1.31 (0.12)***	318,740.5
	Like	−0.85 (0.13)***	195,088.5	0.94 (0.11)***	302,318.0

*** $P < 0.001$: Mann–Whitney *u* tests were conducted between the control group and each of the bias groups.

voters throughout India who had not yet voted (recruited through print advertisements, online advertisements, and online subject pools) to one of three groups in which search rankings favored either Rahul Gandhi, Arvind Kejriwal, or Narendra Modi, the three major candidates in the election.[§]

Subjects were incentivized to participate in the study either with payments between USD\$1 and USD\$4 or with the promise that a donation of approximately USD\$1.50 would be made to a prominent Indian charity that provides free lunches for Indian children. (At the close of the study, a donation of USD\$1,457 was made to the Akshaya Patra Foundation.)

Regarding ethical concerns, because we recruited only a small number of subjects relative to the size of the Indian voting population, we were not concerned that our manipulation could affect the election's outcome. Moreover, our study was designed so that it did not favor any one candidate, so there was no overall bias. The study presented no more than minimal risk to subjects and was approved by AIBRT's IRB. Informed consent was obtained from all subjects.

The subjects ($n = 2,150$) were demographically diverse (Table S5), residing in 27 of 35 Indian states and union territories, and political leanings varied as follows: 13.3% identified themselves as politically right (conservative), 43.8% as center (moderate), 26.0% as left (liberal), and 16.9% as indifferent. In contrast to studies 1 and 2, subjects reported high familiarity with the political candidates (mean familiarity Gandhi, 7.9; SD = 2.5; mean familiarity Kejriwal, 7.7; SD = 2.5; mean familiarity Modi, 8.5; SD = 2.1). The full dataset for all five experiments is accessible at Dataset S1.

Subjects reported more frequent search engine use compared with subjects in studies 1 or 2 (mean searches per day, 15.7; SD = 30.1), and 71.7% of subjects reported that they had previously used a search engine to learn about political candidates. Subjects also spent less time using our mock search engine (mean total time, 277.4 s; SD = 368.3) than did subjects in studies 1 or 2. The patterns of search result clicks and time spent on Web pages in our mock search engine was similar to the patterns we found in study 1 [clicks, $r(28) = 0.96$; $P < 0.001$; Web page time, $r(28) = 0.91$; $P < 0.001$] and study 2 [clicks, $r(28) = 0.96$; $P < 0.001$; Web page time, $r(28) = 0.92$; $P < 0.001$].

Before Web research, we found one significant difference among the three groups for a rating pertaining to Kejriwal, but none for Gandhi or Modi (Table S6). Following the Web research, most of the subjects' ratings of the candidates shifted in the predicted directions (Table 4).

Before Web research, we found no significant differences among the three groups with respect to voting proportions (Table 5). Following Web research, significant differences emerged among the three groups for this measure (Table 5), and the VMP was 10.6% (95% CI, 8.3–12.8%; McNemar's test, $P < 0.001$). Using poststratification and weights obtained from the 2011 India Census data on literate Indians (53)—scaled to size for age, sex, and location (grouped into state or union territory)—the VMP was 9.4% (95% CI, 8.2–10.6%; McNemar's test, $P < 0.001$). When weighted using the same demographics via classical regression post-

stratification (Table S7), the VMP was 9.5% (95% CI, 8.3–10.7%; McNemar's test, $P < 0.001$).

To obtain a more sensitive measure of voting preference in study 3, we asked subjects to estimate the likelihood, on three separate 11-point scales from −5 to +5, that they would vote for each of the candidates if the election were held today. Before Web research, we found no significant differences among the three groups with respect to the likelihood of voting for any of the candidates (Table S6). Following Web research, significant differences emerged among the three groups with respect to the likelihood of voting for Rahul Gandhi and Arvind Kejriwal but not Narendra Modi (Table S6), and all likelihoods shifted in the predicted directions (Table 4). The proportion of subjects showing no awareness of the manipulation in experiment 5 was 99.5%.

In study 3, as in study 2, we found substantial differences in how vulnerable different demographic groups were to SEME (SI Text). Consistent with the findings of study 2 and previous findings on the moderators of order effects (30–32), for example, we found that subjects reporting a low familiarity with the candidates (familiarity less than 5 on a scale from 1 to 10) were more vulnerable to SEME (VMP = 13.7%; 95% CI, 4.3–23.2%; McNemar's test, $P = 0.17$) than were subjects who reported high familiarity with the candidates (VMP = 10.3%; 95% CI, 8.0–12.6%; McNemar's test, $P < 0.001$), although this difference was not significant ($\chi^2 = 0.575$, $P = 0.45$).

As in study 2, although the demographic groups we examined were overlapping and somewhat arbitrary, if one was manipulating an election, information about such differences would have enormous practical value. For example, we found that subjects between ages 18 and 24 were less vulnerable to SEME (VMP = 8.9%; 95% CI, 5.0–12.8%; McNemar's test, $P < 0.05$) than were subjects between ages 45 and 64 (VMP = 18.9%; 95% CI, 6.3–31.5%; McNemar's test, $P = 0.10$) and that self-labeled Christians were more vulnerable (VMP = 30.7%; 95% CI, 20.2–41.1%; McNemar's test, $P < 0.001$) than self-labeled Hindus (VMP = 8.7%; 95% CI, 6.3–11.1%; McNemar's test, $P < 0.001$). Among the most vulnerable groups we identified were unemployed males from Kerala (VMP = 72.7%; 95% CI, 46.4–99.0%; McNemar's test, $P < 0.05$), whereas among the least vulnerable groups were female conservatives (VMP = −11.8%; 95% CI, −29.0%–5.5%; McNemar's test, $P = 0.62$).

A negative VMP might suggest oppositional attitudes or an underdog effect for that group (54). No negative VMPs were found in the demographic groups examined in study 2, but it is understandable that they would be found in an election in which people are highly familiar with the candidates (study 3). As a practical matter, where a search engine company has the ability to send people customized rankings and where biased search rankings are likely to produce an oppositional response with certain voters, such rankings would probably not be sent to them. Eliminating the 2.6% of our sample ($n = 56$) with oppositional responses, the overall VMP in this experiment increases from 10.6% to 19.8% (95% CI, 16.8–22.8%; $n = 2,094$; McNemar's test: $P < 0.001$).

As we found in study 2, awareness of the manipulation appeared to enhance the effect rather than nullify it. The VMP for people

[§]English is one of India's two official languages, the other being Hindi.

Table 4. Postsearch shifts in voting preferences for study 3

Candidate	Rating	χ^2	Mean (SE)		
			Gandhi bias	Kejriwal bias	Modi bias
Gandhi	Impression	3.61	−0.16 (0.06)	−0.21 (0.06)	−0.30 (0.06)
	Trust	21.19***	0.14 (0.06)	−0.04 (0.07)	−0.20 (0.06)
	Like	12.99**	−0.09 (0.07)	−0.17 (0.06)	−0.34 (0.06)
	Voting likelihood	10.79**	0.16 (0.07)	−0.04 (0.07)	−0.18 (0.07)
Kejriwal	Impression	17.75***	−0.30 (0.06)	−0.11 (0.06)	−0.39 (0.05)
	Trust	26.69***	−0.17 (0.07)	0.15 (0.06)	−0.16 (0.06)
	Like	24.74***	−0.31 (0.06)	0.05 (0.06)	−0.23 (0.06)
	Voting likelihood	13.22**	−0.03 (0.06)	0.17 (0.07)	−0.12 (0.06)
Modi	Impression	24.98***	−0.22 (0.06)	−0.21 (0.06)	0.12 (0.05)
	Trust	18.78***	−0.04 (0.06)	−0.10 (0.06)	0.23 (0.06)
	Like	16.89***	−0.16 (0.05)	−0.09 (0.06)	0.19 (0.06)
	Voting likelihood	31.07***	−0.07 (0.07)	−0.10 (0.06)	0.33 (0.06)

** $P < 0.01$ and *** $P < 0.001$: for each rating, a Kruskal–Wallis χ^2 test was used to assess significance of group differences.

who showed no awareness of the biased search rankings ($n = 2,140$) was 10.5% (95% CI, 8.3–12.7%; McNemar's test, $P < 0.001$), whereas the VMP for people who showed awareness of the bias ($n = 10$) was 33.3%.

The rankings and Web pages we used in study 3 were selected by the investigators based on our limited understanding of Indian politics and perspectives. To optimize the rankings, midway through the election process we hired a native consultant who was familiar with the issues and perspectives pertinent to undecided voters in the 2014 Lok Sabha Election. Based on the recommendations of the consultant, we made slight changes to our rankings on 30 April, 2014. In the preoptimized rankings group ($n = 1,259$), the VMP was 9.5% (95% CI, 6.8–12.2%; McNemar's test, $P < 0.001$); in the postoptimized rankings group ($n = 891$), the VMP increased to 12.3% (95% CI, 8.5–16.1%; McNemar's test, $P < 0.001$). Eliminating the 3.1% of the subjects in the postoptimization sample with oppositional responses ($n = 28$), the VMP increased to 24.5% (95% CI, 19.3–29.8%; $n = 863$).

Discussion

Elections are often won by small vote margins. Fifty percent of US presidential elections were won by vote margins under 7.6%, and 25% of US senatorial elections in 2012 were won by vote margins under 6.0% (55, 56). In close elections, undecided voters can make all of the difference, which is why enormous resources are often focused on those voters in the days before the election (57, 58). Because search rankings biased toward one candidate can apparently sway the voting preferences of undecided voters without their awareness and, at least under some circumstances, without any possible competition from opposing candidates, SEME appears to be an especially powerful tool for manipulating elections. The Australian election used in studies 1 and 2 was won by a margin of only 0.24% and perhaps could easily have been turned by such a manipulation. The Fox News Effect, which is small compared with SEME, is believed to have shifted between 0.4% and 0.7% of votes to conservative candidates:

more than enough, according to the researchers, to have had a “decisive” effect on a number of close elections in 2000 (40).

Political scientists have identified two of the most common methods political candidates use to try to win elections. The core voter model describes a strategy in which resources are devoted to mobilizing supporters to vote (59). As noted earlier, Zittrain recently pointed out that a company such as Facebook could mobilize core voters to vote on election day by sending “get-out-and-vote” messages en masse to supporters of only one candidate. Such a manipulation could be used undetectably to flip an election in what might be considered a sort of digital gerrymandering (44, 45). In contrast, the swing voter model describes a strategy in which candidates target their resources toward persuasion—attempting to change the voting preferences of undecided voters (60). SEME is an ideal method for influencing such voters.

Although relatively few voters have actively sought political information about candidates in the past (61), the ease of obtaining information over the Internet appears to be changing that: 73% of online adults used the Internet for campaign-related purposes during the 2010 US midterm elections (61), and 55% of all registered voters went online to watch videos related to the 2012 US election campaign (62). Moreover, 84% of registered voters in the United States were Internet users in 2012 (62). In our nationwide study in the United States (study 2), 86.0% of our subjects reported having used search engines to get information about candidates. Meanwhile, the number of people worldwide with Internet access is increasing rapidly, predicted to increase to nearly 4 billion by 2018 (63). By 2018, Internet access in India is expected to rise from the 213 million users who had access in 2013 to 526 million (63). Worldwide, it is reasonable to conjecture that both proportions will increase substantially in future years; that is, more people will have Internet access, and more people will obtain information about candidates from the Internet. In the context of the experiments we have presented, this suggests that whatever the effect sizes we have observed now, they will likely be larger in the future.

Table 5. Comparison of voting proportions before and after Web research for study 3

Group	Simulated vote before Web research			χ^2	Simulated vote after Web research			χ^2	VMP
	Gandhi	Kejriwal	Modi		Gandhi	Kejriwal	Modi		
1	115	164	430	3.070	144	152	413	16.935**	10.6%***
2	112	183	393		113	199	376		
3	127	196	430		117	174	462		

McNemar's test was conducted to assess VMP significance. VMP, percent increase in subjects in the bias groups combined who said that they would vote for the favored candidate.

** $P < 0.01$; and *** $P < 0.001$: Pearson χ^2 tests were conducted among all three groups.

The power of SEME to affect elections in a two-person race can be roughly estimated by making a small number of fairly conservative assumptions. Where i is the proportion of voters with Internet access, u is the proportion of those voters who are undecided, and VMP, as noted above, is the proportion of those undecided voters who can be swayed by SEME, W —the maximum win margin controllable by SEME—can be estimated by the following formula: $W = i \cdot u \cdot \text{VMP}$.

In a three-person race, W will vary between 75% and 100% of its value in a two-person race, depending on how the votes are distributed between the two losing candidates. (Derivations of formulas in the two-candidate and three-candidate cases are available in *SI Text*.) In both cases, the size of the population is irrelevant.

Knowing the values for i and u for a given election, along with the projected win margin, the minimum VMP needed to put one candidate ahead can be calculated (Table S8). In theory, continuous online polling would allow search rankings to be optimized continuously to increase the value of VMP until, in some instances, it could conceivably guarantee an election's outcome, much as "conversion" and "click-through" rates are now optimized continuously in Internet marketing (64).

For example, if (i) 80% of eligible voters had Internet access, (ii) 10% of those individuals were undecided at some point, and (iii) SEME could be used to increase the number of people in the undecided group who were inclined to vote for the target candidate by 25%, that would be enough to control the outcome of an election in which the expected win margin was as high as 2%. If SEME were applied strategically and repeatedly over a period of weeks or months to increase the VMP, and if, in some locales and situations, i and u were larger than in the example given, the controllable win margin would be larger. That possibility notwithstanding, because nearly 25% of national elections worldwide are typically won by margins under 3%,[†] SEME could conceivably impact a substantial number of elections today even with fairly low values of i , u , and VMP.

Given our procedures, however, we cannot rule out the possibility that SEME produces only a transient effect, which would limit its value in election manipulation. Laboratory manipulations of preferences and attitudes often impact subjects for only a short time, sometimes just hours (65). That said, if search rankings were being manipulated with the intent of altering the outcome of a real election, people would presumably be exposed to biased rankings repeatedly over a period of weeks or months. We produced substantial changes not only in voting preferences but in multiple ratings of attitudes toward candidates given just one exposure to search rankings linking to Web pages favoring one candidate, with average search times in the 277- to 635-s range. Given hundreds or thousands of exposures of this sort, we speculate not only that the resulting attitudes and preferences would be stable, but that they would become stronger over time, much as brand preferences become stronger when advertisements are presented repeatedly (66).

Our results also suggest that it is a relatively simple matter to mask the bias in search rankings so that it is undetectable to virtually every user. In experiment 3, using only a simple mask, none of our subjects appeared to be aware that they were seeing biased rankings, and in our India study, only 0.5% of our subjects appeared to notice the bias. When people are subjected to forms of influence they can identify—in campaigns, that means speeches, billboards, television commercials, and so on—they can defend themselves fairly easily if they have opposing views. Invisible sources of influence can be harder to defend against (67–69), and for people who are impressionable, invisible sources of influence not only persuade, they also leave people feeling that they made up their own minds—that no external force was applied (70, 71). Influence is sometimes undetectable because key stimuli act subliminally (72–74), but

search results and Web pages are easy to perceive; it is the pattern of rankings that people cannot see. This invisibility makes SEME especially dangerous as a means of control, not just of voting behavior but perhaps of a wide variety of attitudes, beliefs, and behavior. Ironically, and consistent with the findings of other researchers, we found that even those subjects who showed awareness of the biased rankings were still impacted by them in the predicted directions (75).

One weakness in our studies was the manner in which we chose to determine whether subjects were aware of bias in the search rankings. As noted, to not generate false-positive responses, we avoided asking leading questions that referred specifically to bias; rather, we asked a rather vague question about whether anything had bothered subjects about the search rankings, and we then gave subjects an opportunity to type out the details of their concerns. In so doing, we probably underestimated the number of detections (47), and this is a matter that should be studied further. That said, because people who showed awareness of the bias were still vulnerable to our manipulation, people who use SEME to manipulate real elections might not be concerned about detection, except, perhaps, by regulators.

Could regulators in fact detect SEME? Theoretically, by rating pages and monitoring search rankings on an ongoing basis, search ranking bias related to elections might be possible to identify and track; as a practical matter, however, we believe that biased rankings would be impossible or nearly impossible for regulators to detect. The results of studies 2 and 3 suggest that vulnerability to SEME can vary dramatically from one demographic group to another. It follows that if one were using biased search rankings to manipulate a real election, one would focus on the most vulnerable demographic groups. Indeed, if one had access to detailed online profiles of millions of individuals, which search engine companies do (76–78), one would presumably be able to identify those voters who appeared to be undecided and impressionable and focus one's efforts on those individuals only—a strategy that has long been standard in political campaigns (79–84) and continues to remain important today (85). With search engine companies becoming increasingly adept at sending users customized search rankings (76–78, 86–88), it seems likely that only customized rankings would be used to influence elections, thus making it difficult or impossible for regulators to detect a manipulation. Rankings that appear to be unbiased on the regulators' screens might be highly biased on the screens of select individuals.

Even if a statistical analysis did show that rankings consistently favored one candidate over another, those rankings could always be attributed to algorithm-guided dynamics driven by market forces—so-called "organic" forces (89)—rather than by deliberate manipulation by search engine company employees. This possibility suggests yet another potential danger of SEME. What if election-related search results are indeed being left to the vagaries of market forces? Do such forces end up pushing some candidates to the top of search rankings? If so, it seems likely that those high rankings are cultivating additional supporters for those candidates in a kind of digital bandwagon effect. In other words, for several years now and with greater impact each year (as more people get election-related information through the Internet), SEME has perhaps already been affecting the outcomes of close elections. To put this another way, without human intention or direction, algorithms have perhaps been having a say in selecting our leaders.

Because search rankings are based, at least in part, on the popularity of Web sites (90), it is likely that voter preferences impact those rankings to some extent. Given our findings that search rankings can in turn affect voter preferences, these phenomena might interact synergistically, causing a substantial increase in support for one candidate at some point even when the effects of the individual phenomena are small.[‡]

Our studies produced a wide range of VMPs. In a real election, what proportion of undecided voters could actually be

[†]Some of the data applied in the analysis in this publication are based on material from the "European Election Database." The data are collected from original sources and prepared and made available by the Norwegian Social Science Data Services (NSD). NSD is not responsible for the analyses/interpretation of the data presented here.

[‡]A mathematical model we developed—highly conjectural, we admit, and at this point unverifiable—shows the possible dynamics of such synergy (Fig. S1).

shifted using SEME? Our first two studies, which relied on a campaign and candidates that were unfamiliar to our subjects, produced overall VMPs in the range 36.7–63.3%, with demographic shifts occurring with VMPs as high as 80.0%. Our third study, with real voters in the midst of a real election, produced, overall, a lower VMP: just 10.6%, with optimizing our rankings raising the VMP to 12.3% and with the elimination of a small number of oppositional subjects raising the VMP to 24.5%, which is the value we would presumably have found if our search rankings had been optimized from the start and if we had advance knowledge about oppositional groups. In the third study, VMPs in some demographic groups were as high as 72.7%. If a search engine company optimized rankings continuously and sent customized rankings only to vulnerable undecided voters, there is no telling how high the VMP could be pushed, but it would almost certainly be higher than our modest efforts could achieve. Our investigation suggests that with optimized, targeted rankings, a VMP of at least 20% should be relatively easy to achieve in real elections. Even if only 60% of a population had Internet access and only 10% of voters were undecided, that would still allow control of elections with win margins up to 1.2%—five times greater than the win margin in the 2010 race between Gillard and Abbott in Australia.

Conclusions

Given that search engine companies are currently unregulated, our results could be viewed as a cause for concern, suggesting that such companies could affect—and perhaps are already affecting—the outcomes of close elections worldwide. Restricting search ranking manipulations to voters who have been identified as undecided while also donating money to favored candidates would be an especially subtle, effective, and efficient way of wielding influence.

Although voters are subjected to a wide variety of influences during political campaigns, we believe that the manipulation of search rankings might exert a disproportionately large influence over voters for four reasons:

First, as we noted, the process by which search rankings affect voter preferences might interact synergistically with the process by which voter preferences affect search rankings, thus creating a sort of digital bandwagon effect that magnifies the potential impact of even minor search ranking manipulations.

Second, campaign influence is usually explicit, but search ranking manipulations are not. Such manipulations are difficult

to detect, and most people are relatively powerless when trying to resist sources of influence they cannot see (66–68). Of greater concern in the present context, when people are unaware they are being manipulated, they tend to believe they have adopted their new thinking voluntarily (69, 70).

Third, candidates normally have equal access to voters, but this need not be the case with search engine manipulations. Because the majority of people in most democracies use a search engine provided by just one company, if that company chose to manipulate rankings to favor particular candidates or parties, opponents would have no way to counteract those manipulations. Perhaps worse still, if that company left election-related search rankings to market forces, the search algorithm itself might determine the outcomes of many close elections.

Finally, with the attention of voters shifting rapidly toward the Internet and away from traditional sources of information (12, 61, 62), the potential impact of search engine rankings on voter preferences will inevitably grow over time, as will the influence of people who have the power to control such rankings.

We conjecture, therefore, that unregulated election-related search rankings could pose a significant threat to the democratic system of government.

Materials and Methods

We used 102 subjects in each of experiments 1–3 to give us an equal number of subjects in all three groups and both counterbalancing conditions of the experiments.

Nonparametric statistical tests such as the Mann–Whitney u and the Kruskal–Wallis H are used throughout the present report because Likert scale scores, which were used in each of the studies, are ordinal.

In study 3, the procedure was identical to that of studies 1 and 2; only the Web pages and search results were different: that is, Web pages and search results were pertinent to the three leading candidates in the 2014 Lok Sabha general elections. The questions we asked subjects were also adjusted for a three-person race.

ACKNOWLEDGMENTS. We thank J. Arnett, E. Clemons, E. Fantino, S. Glenn, M. Hovell, E. Key, E. Loftus, C. McKenzie, B. Meredith, N. Metaxas, D. Moriarty, D. Peel, M. Runco, S. Stolarz-Fantino, and J. Wixted for comments; K. Robertson for image editing; V. Sharan for advice on optimizing search rankings in the India study; K. Duncan and F. Tran for assistance with data analysis; J. Hagan for technical assistance; and S. Palacios and K. Huynh for assistance in conducting the experiments in study 1. This work was supported by the American Institute for Behavioral Research and Technology, a nonpartisan, nonprofit organization.

- Joachims T, et al. (2007) Evaluating the accuracy of implicit feedback from clicks and query reformulations in web search. *Association for Computing Machinery Transactions on Information Systems* 25(2):7.
- Pan B, et al. (2007) In Google we trust: User's decisions on rank, position, and relevance. *J Comput Mediat Commun* 12(3):801–823.
- Guan Z, Cutrell E (2007) An eye tracking study of the effect of target rank on web search. *Proceedings of the Special Interest Group for Computer-Human Interaction Conference on Human Factors in Computing Systems* (ACM, New York), pp 417–420.
- Lorigo L, et al. (2008) Eye tracking and online search: Lessons learned and challenges ahead. *J Am Soc Inf Sci Technol* 59(7):1041–1052.
- Granka LA, Joachim T, Gay G (2004) Eye-tracking analysis of user behavior in www search. *Proceedings of the 27th Annual International Association for Computing Machinery Special Interest Group on Information Retrieval Conference on Research and Development in Information Retrieval* (ACM, New York, NY), pp 478–479.
- Agichtein E, Brill E, Dumais S, Ragno R (2006) Learning user interaction models for predicting web search result preferences. *Proceedings of the 29th Annual International Association for Computing Machinery Special Interest Group on Information Retrieval Conference on Research and Development in Information Retrieval* (ACM, New York, NY), pp 3–10.
- Chitika (2013) The value of Google result positioning. Available at perma.cc/7AGC-HTDH. Accessed June 30, 2015.
- Optify (2011) The changing face of SERPS: Organic CTR. Available at perma.cc/KZ5X-78TL. Accessed June 30, 2015.
- Jansen BJ, Spink A, Saracevic T (2000) Real life, real users, and real needs: A study and analysis of user queries on the web. *Inf Process Manage* 36(2):207–227.
- Spink A, Wolfram D, Jansen BJ, Saracevic T (2001) Searching the web: The public and their queries. *J Am Soc Inf Sci Technol* 53(3):226–234.
- Silverstein C, Marais H, Henzinger M, Moricz A (1999) Analysis of a very large web search engine query log. *Association for Computing Machinery Special Interest Group on Information Retrieval Forum* 33(1):6–12.
- Purcell K, Brenner J, Rainie L (2012) Pew Research Center's Internet & American Life project: Search engine use 2012. Available at perma.cc/NF6C-JCPW. Accessed June 30, 2015.
- Learmonth M (2010) What big brands are spending on Google. Available at perma.cc/5L3B-SPTX. Accessed June 30, 2015.
- Econsultancy (2012) SEMPO state of search marketing report 2012. Available at perma.cc/A3KH-6QVB. Accessed June 30, 2015.
- Gerhart S (2004) Do Web search engines suppress controversy? *First Monday* 9(1):1111.
- Ebbinghaus H (1913) *Memory: A Contribution to Experimental Psychology* (No. 3) (Teachers College, Columbia Univ, New York).
- Murdock BB (1962) The serial position effect of free recall. *J Exp Psychol* 64(5):482–488.
- Asch SE (1946) Forming impressions of personality. *J Abnorm Psychol* 41:258–290.
- Lund FH (1925) The psychology of belief. *J Abnorm Soc Psych* 20(1):63–81.
- Hovland CI (1957) *The Order of Presentation in Persuasion* (Yale Univ Press, New Haven, CT).
- Baird JE, Zelin RC (2000) The effects of information ordering on investor perceptions: An experiment utilizing presidents' letters. *J Financial Strategic Decisions* 13(3):71–80.
- Tourangeau R, Couper MP, Conrad FG (2013) "Up means good" The effect of screen position on evaluative ratings in web surveys. *Public Opin Q* 77(1, Suppl 1):69–88.
- Krosnick JA, Alwin DF (1987) An evaluation of a cognitive theory of response order effects in survey measurement. *Pub Op Quarterly* 51(2):201–219.
- Stern MJ, Dillman DA, Smyth JD (2007) Visual design, order effects, and respondent characteristics in a self-administered survey. *Surv Res Methods* 1(3):121–138.
- Miller JE (1980) *Menu Pricing and Strategy* (CBI Publishing, Boston).
- Ho DE, Imai K (2008) Estimating causal effects of ballot order from a randomized natural experiment the California alphabet lottery, 1978–2002. *Pub Op Quarterly* 72(2):216–240.
- Chen E, Simonovits G, Krosnick JA, Pasek J (2014) The impact of candidate name order on election outcomes in North Dakota. *Elect Stud* 35:115–122.

28. Krosnick JA, Miller JM, Tichy MP (2004) An unrecognized need for ballot reform. *Rethinking the Vote*, eds Crigler A, Just M, McCaffery E (Oxford Univ Press, New York), pp 51–74.
29. Koppell JG, Steen JA (2004) The effects of ballot position on election outcomes. *J Polit* 66(1):267–281.
30. Kim N, Krosnick JA, Casasanto D (2014) Moderators of candidate name-order effects in elections: An experiment. *Polit Psychol*, 10.1111/pops.12178.
31. Miller JM, Krosnick JA (1998) The impact of candidate name order on election outcomes. *Pub Op Quarterly* 62(3):291–330.
32. Pasek J, et al. (2014) Prevalence and moderators of the candidate name-order effect evidence from statewide general elections in California. *Pub Op Quarterly* 78(2):416–439.
33. Murphy J, Hofacker C, Mizerski R (2006) Primacy and recency effects on clicking behavior. *J Comput Mediat Commun* 11(2):522–535.
34. Ansari A, Mela CF (2003) E-customization. *J Mark Res* 40(2):131–145.
35. Drèze X, Zufryden F (2004) Measurement of online visibility and its impact on Internet traffic. *J Interact Market* 18:20–37.
36. Chiang CF, Knight BG (2011) Media bias and influence: Evidence from newspaper endorsements. *Rev Econ Stud* 78(3):795–820.
37. Druckman JN, Parkin M (2005) The impact of media bias: How editorial slant affects voters. *J Polit* 67(4):1030–1049.
38. Gerber A, Karlan DS, Bergan D (2009) Does the media matter? A field experiment measuring the effect of newspapers on voting behavior and political opinions. *Am Econ J Appl Econ* 1(2):35–52.
39. Morwitz VG, Pluzinski C (1996) Do pools reflect opinions or do opinions reflect polls? The impact of political polling on voters' expectation, preferences, and behavior. *J Consum Res* 23(1):53–67.
40. DellaVigna S, Kaplan E (2007) The Fox News effect: Media bias and voting. *Q J Econ* 122(3):1187–1234.
41. Lakoff G (1987) *Women, Fire and Dangerous Things* (Univ of Chicago Press, Chicago).
42. DeMarzo PM, Vayanos D, Zwiabell J (2003) Persuasion bias, social influence, and multidimensional opinions. *Q J Econ* 118(3):909–968.
43. Fogg BJ (2002) *Persuasive Technology: Using Computers to Change What We Think and Do* (Morgan Kaufmann, San Francisco).
44. Bond RM, et al. (2012) A 61-million-person experiment in social influence and political mobilization. *Nature* 489(7415):295–298.
45. Zittrain J (2014) Engineering an election. Digital gerrymandering poses a threat to democracy. *Harvard Law Review Forum* 127(8):335–341.
46. US Census Bureau (2011) Voting and registration in the election of November 2010. Available at perma.cc/DU9N-C6XT. Accessed June 30, 2015.
47. Loftus EF (1996) *Eyewitness Testimony* (Harvard Univ Press, Boston).
48. Kühberger A, Fritz A, Scherndl T (2014) Publication bias in psychology: A diagnosis based on the correlation between effect size and sample size. *PLoS One* 9(9):e105825.
49. Paolacci G, Chandler J (2014) Inside the Turk: Understanding mechanical Turk as a participant pool. *Curr Dir Psychol Sci* 23(3):184–188.
50. Berinsky AJ, Huber GA, Lenz GS (2012) Evaluating online labor markets for experimental research: Amazon.com's mechanical turk. *Polit Anal* 20(3):351–368.
51. Gates GJ (2011) How many people are lesbian, gay, bisexual and transgender? Available at perma.cc/N7P2-TXHV. Accessed June 30, 2015.
52. Gelman A, Hill J (2006) *Data Analysis Using Regression and Multilevel/Hierarchical Models* (Cambridge Univ Press, Boston).
53. Census of India (2011) Population enumeration data. Available at perma.cc/4JKF-UWQZ. Accessed June 30, 2015.
54. Simon HA (1954) Bandwagon and underdog effects and the possibility of election predictions. *Pub Op Quarterly* 18(3):245–253.
55. Leip D (2012) Dave Leip's atlas of U.S. presidential elections. Available at perma.cc/G9YF-6CPE. Accessed June 30, 2015.
56. Rogers D, Cage F (2012) US house and senate elections 2012. Available at perma.cc/5QSK-FNCL. Accessed June 30, 2015.
57. Cabletelevision Advertising Bureau (2012) CAB undecided voters study 2012. Available at perma.cc/T82K-YC38. Accessed June 30, 2015.
58. Rutenberg J (2012) Obama campaign endgame: Grunt work and cold math. Available at perma.cc/X74D-YNM5. Accessed June 30, 2015.
59. Cox GW, McCubbins MD (1986) Electoral politics as a redistributive game. *J Polit* 48(2):370–389.
60. Lindbeck A, Weibull JW (1987) Balanced-budget redistribution as the outcome of political competition. *Public Choice* 52(3):273–297.
61. Smith A (2011) Pew Research Center's Internet & American Life Project: The Internet and campaign 2010. Available at perma.cc/3U9R-EUQQ. Accessed June 30, 2015.
62. Smith A, Duggan M (2012) Pew Research Center's Internet & American life project: Online political videos and campaign 2012. Available at perma.cc/5UX2-LELJ. Accessed June 30, 2015.
63. Cisco (2014) VNI Global IP Traffic Forecast, 2013–2018. Available at perma.cc/AE9R-JSHL. Accessed June 30, 2015.
64. Ash T, Ginty M, Page R (2012) *Landing Page Optimization: The Definitive Guide to Testing and Tuning for Conversions* (John Wiley & Sons, Indianapolis).
65. Kraus SJ (1995) Attitudes and the prediction of behavior: A meta-analysis of the empirical literature. *Pers Soc Psychol Bull* 21(1):58–75.
66. Zajonc RB (2001) Mere exposure: A gateway to the subliminal. *Curr Dir Psychol Sci* 10(6):224–228.
67. Berger J, Fitzsimons G (2008) Dogs on the street, pumas on your feet: How cues in the environment influence product evaluation and choice. *J Mark Res* 45(1):1–14.
68. Brasel SA, Gips J (2011) Red bull “gives you wings” for better or worse: A double-edged impact of brand exposure on consumer performance. *J Consum Psychol* 21(1):57–64.
69. Fransen ML, Fennis BM, Pruyn ATH, Das E (2008) Rest in peace? Brand-induced mortality salience and consumer behavior. *J Bus Res* 64(10):1053–1061.
70. Bargh JA, Gollwitzer PM, Lee-Chai A, Barndollar K, Trötschel R (2001) The automated will: Nonconscious activation and pursuit of behavioral goals. *J Pers Soc Psychol* 81(6):1014–1027.
71. Pronin E, Kugler MB (2007) Valuing thoughts, ignoring behavior: The introspection illusion as a source of the bias blind spot. *J Exp Soc Psychol* 43(4):565–578.
72. Winkielman P, Berridge KC, Wilbarger JL (2005) Unconscious affective reactions to masked happy versus angry faces influence consumption behavior and judgments of value. *Pers Soc Psychol Bull* 31(1):121–135.
73. Karremans JC, Stroebe W, Claus J (2006) Beyond Vicary's fantasies: The impact of subliminal priming and brand choice. *J Exp Soc Psychol* 42(6):792–798.
74. Légal JB, Chappé J, Coiffard V, Villard-Forest A (2012) Don't you know that you want to trust me? Subliminal goal priming and persuasion. *J Exp Soc Psychol* 48(1):358–360.
75. Mills J, Aronson E (1965) Opinion change as a function of the communicator's attractiveness and desire to influence. *J Pers Soc Psychol* 1(2):173–177.
76. Bi B, Shokouhi M, Kosinski M, Graepel T (2013) Inferring the demographics of search users: Social data meets search queries. *Proceedings of the 22nd International Conference on World Wide Web* (International World Wide Web Conferences Steering Committee, Republic and Canton of Geneva, Switzerland), pp 131–140.
77. Yan J, et al. (2009, April). How much can behavioral targeting help online advertising? *Proceedings of the 18th International Conference on World Wide Web* (ACM, New York), pp 261–270.
78. Weber I, Castillo C (2010) The demographics of web search. *Proceedings of the 33rd Annual International Association for Computing Machinery Special Interest Group on Information Retrieval Conference on Research and Development in Information Retrieval* (ACM, New York), pp 523–530.
79. Mayer WG (2008) *The Swing Voter in American Politics* (Brookings Institution Press, Washington, DC).
80. Cox GW (2010) Swing voters, core voters, and distributive politics. *Political Representation*, eds Shapiro I, Stokes SC, Wood EJ, Kirshner AS (Cambridge Univ Press, Cambridge, MA), pp 342–357.
81. Bickers KN, Stein RM (1996) The electoral dynamics of the federal pork barrel. *Am J Pol Sci* 40(4):1300–1325.
82. Dahlberg M, Johansson E (2002) On the vote-purchasing behavior of incumbent governments. *Am Polit Sci Rev* 96(1):27–40.
83. Stein RM, Bickers KN (1994) Congressional Elections and the Pork Barrel. *J Polit* 56(2):377–399.
84. Stokes SC (2005) Perverse accountability: A formal model of machine politics with evidence from Argentina. *Am Polit Sci Rev* 99(3):315–325.
85. Mayer WG (2012) The disappearing—but still important—swing voter. *Forum* 10(3):1520.
86. Agichtein E, Brill E, Dumais S (2006) Improving web search ranking by incorporating user behavior information. *Proceedings of the 29th Annual International Association for Computing Machinery Special Interest Group on Information Retrieval Conference on Research and Development in Information Retrieval* (ACM, New York), pp 19–26.
87. Ahmed A, Aly M, Das A, Smola AJ, Anastasakos T (2012) Web-scale multi-task feature selection for behavioral targeting. *Proceedings of the 21st Association for Computing Machinery International Conference on Information and Knowledge Management* (ACM, New York), pp 1737–1741.
88. Chen Y, Pavlov D, Canny JF (2009) Large-scale behavioral targeting. *Proceedings of the 15th Association for Computing Machinery Special Interest Group on Knowledge Discovery in Data International Conference on Knowledge Discovery and Data Mining* (ACM, New York), pp 209–218.
89. Höchstötter N, Lewandowski D (2009) What users see—Structures in search engine results pages. *Inf Sci* 179(12):1796–1812.
90. Page L, Brin S, Motwani R, Winograd T (1999) *The PageRank Citation Ranking: Bringing Order to the Web*. Technical Report (Stanford InfoLab, Palo Alto, CA).