

MEDICINE AND SOCIETY

Do Conflict of Interest Disclosures Facilitate Public Trust?

Daylian M. Cain, PhD and Mohin Banker

Abstract

Lab experiments disagree on the efficacy of disclosure as a remedy to conflicts of interest (COIs). Some experiments suggest that disclosure has perverse effects, although others suggest these are mitigated by real-world factors (eg, feedback, sanctions, norms). This article argues that experiments reporting positive effects of disclosure often lack external validity: disclosure works best in lab experiments that make it unrealistically clear that the one disclosing is intentionally lying. We argue that even disclosed COIs remain dangerous in settings such as medicine where bias is often unintentional rather than the result of intentional corruption, and we conclude that disclosure might not be the panacea many seem to take it to be.

Introduction

While most medical professionals have the best intentions, conflicts of interest (COIs) can unintentionally bias their advice.¹ For example, physicians might have consulting relationships with a company whose product they might prescribe. Physicians are increasingly required to limit COIs and disclose any that exist. When regulators decide whether to let a COI stand, the question becomes: How well does disclosure work? This paper reviews laboratory experiments that have had mixed results on the effects of disclosing COIs on bias and suggests that studies purporting to provide evidence of the efficacy of disclosure often lack external validity. We conclude that disclosure works more poorly than regulators hope; thus, COIs are more problematic than expected.

Perverse Effects of Disclosure

Several studies have reported positive effects of disclosure. Koch and Schmidt's recent lab experiments suggest that disclosure reduces bias in advice when audiences receive feedback and advisors can form reputations.²

Similarly, Church and Kuang argue that disclosure mitigates bias when the audience can sanction advisors for giving bad advice.³ Furthermore, Sah argues that disclosure reduces bias in clinical settings because practitioners operate under the ethical norm of “clients first.”⁴ The problem, as we shall explain, is that these experiments rely on disclosures that make it unrealistically clear that advisors are intentionally lying to advisees.

The experiments were a response to earlier studies conducted by Cain, Loewenstein, and Moore (CLM)^{5,6} that suggest disclosure might have perverse effects. For example, CLM argued that disclosure can increase bias in advice due to 2 possible psychological mechanisms. *Moral licensing* to bias advice suggests that, postdisclosure, advisors (perhaps unintentionally) show less self-restraint because “the patient has been warned.”⁵ Prior to disclosure, conflicted advisors rein in their bias; they want to help themselves, but they also (or even primarily) want to help their advisees. Postdisclosure, they might feel less obliged to help their advisees if they think that the advisees can help themselves, having been warned. Second, postdisclosure, advisors might use *strategic exaggeration* to further bias their advice in order to counteract presumed advice-discounting from advisees. It’s as if disclosure causes advisees to cover their ears—and also encourages advisors to yell even louder.

Sah, Loewenstein, and Cain^{7,8} demonstrated further perverse effects, including the *burden of disclosure*,⁷ whereby disclosure causes advisees to feel burdened to follow biased advice. After disclosure, advisees are concerned about the advice being untrustworthy, but they also want to avoid *being seen* as noncompliant⁷ or distrusting of the advisor.⁸ This compliance diminishes if advisees can quietly “exit” the prying eyes of advisors, hide their noncompliance, or somehow make another excuse for noncompliance other than distrust.⁹ A more basic perverse effect is that overreliance on disclosure might supplant efforts to reduce COIs; although this idea is less psychologically complex, it is perhaps the most consequential.

In addition to perverse (backfire) effects, disclosure might simply fall short. For example, regulators often call for more frequent, easier-to-understand disclosures. Although disclosures buried in fine-print legalese help only those doing the burying, research on anchoring and insufficient adjustment¹⁰ suggests that *even when audiences are clearly warned* that the advice was randomly generated, they are still affected by the advice. Thus, disclosures might not totally undo the damage of biased advice, regardless of how clear the disclosures are.

Prodisclosure Research and Its Limitation

Despite these findings on the weaknesses of disclosure, other studies (Koch and Schmidt,² Church and Kuang,³ and Sah⁴) have sought to defend disclosure. However, in these experiments, what is disclosed is clearly identifiable, intentional lying. Lying is often not present in medical contexts—or, at least, not easily identified. For example, consider physicians who had [business relationships](#) with makers of opioids during the overprescription crisis (eg, through taking consulting gigs, abundant “free samples,” or even traditional rewards for treating patients). Even in cases of overprescription, it is likely that many conflicted physicians reasonably—or, at least *plausibly*—believed the drugs were appropriate to alleviate pain. After all, even many of those who advised that Enron was a “strong buy” plausibly believed their recommendations.¹¹ Our point is that if mere disclosure made it easy to prove who was intentionally giving self-interested advice, prodisclosure arguments would unsurprisingly win out. Unfortunately, it is not so easy to identify intentionally biased advice in real-world contexts. And in real-world contexts, COIs often lead to unintentional bias rather than intentional lies.¹²

Granted, even CLM’s own experiments sometimes examined intentionally biased advice. For example, in one study, CLM had advisors rate the ethicality of intentionally providing advice outside a range containing the actual number (of jelly beans in a jar).⁶ However, in the *main* CLM experiments, advisors were asked to give advice that was within a broad range of plausible values⁵ or else no range was given.⁶ Whether or not advisors’ bias was intentional, it was realistically difficult for advisees to know if advisors believed the advice. In other words, CLM’s advisors had *plausible deniability*. Research has shown plausible deniability to be crucial to advisors, even in one-shot experiments.¹³ It is easy to imagine why plausible deniability would be important in the real world—not only to intentional liars who seek protection from litigation, but also to the unintentionally biased who could not otherwise escape (perhaps their own) scrutiny.

Similar to CLM,^{5,6} Koch and Schmidt² tested how advisors’ disclosure of COIs affected the advice they gave when they knew the range of true values. In both CLM’s and Koch and Schmidt’s studies, advisors gave numerical advice to advisees playing numerical guessing games (eg, guessing a random value, estimating the value of coin jars, estimating sale prices of local houses). The advisors had COIs because they were paid more when advisees overestimated the value of the item in question. Koch and Schmidt provided very narrow ranges of the true value to the advisors, and many of their

advisors gave (knowingly false) advice that was outside this range.² Conversely, CLM gave advisors less information about the true value (broader ranges), so CLM's advisors could plausibly deny giving bad advice if it remained within the given range of values. The studies incorporated feedback of advisors' and advisees' estimated values that could be taken into account in the next round of advising; however, in Koch and Schmidt's study, advisee feedback often made it unrealistically clear that the advisor was lying because their estimates were outside the range of true values. As a result, disclosure that the advisor had a COI would be especially damning when coupled with the now obvious fact that the advisor had lied in the prior round. It is not realistic for advisees to receive such detailed external feedback or for advisors to even know the range of true values. The advisors in CLM's studies disclosed COIs but often could have plausibly given well-intentioned advice because the range of true values was so broad. The difference is one of being warned that *your physician intentionally lies to you* vs being warned that *your physician might be biased*.

Similar problems abound in Church and Kuang's study on combining disclosure with sanctions.³ Advisors knew that advice outside a certain range would be unequivocally wrong, but the findings suggest that many advisors still gave intentionally wrong advice (ie, outside the true range) when COIs were not disclosed. Disclosure would highlight the possibility that advisors were lying or biased, so it is not surprising that advisors would lie less when liars could easily be punished: advisees merely needed to select sanctioning options, and liars were automatically punished by the experimental system, regardless of whether advisees were aware that the advisor was lying. Church and Kuang admit to this limitation, stating, "In our setting, an adviser who provided bad advice would be penalized with certainty, as long as the investor chose to initiate sanctions."³ They credit an anonymous reviewer for pointing out the problem here: "in many naturally occurring settings, when advice turns out to be bad, it might be difficult to discern whether that is due to the adviser's bias or uncontrollable factors such as environmental volatility. As a result, biased advisers [Cain and Banker would add: 'in the real-world'] are not necessarily penalized.... We acknowledge that under such circumstances, the investor's threat of initiating sanctions might have less teeth than in our setting." At least Church and Kuang acknowledge this limitation: disclosure reduces bias when sanctions have (unrealistically) sharp teeth and bad advisors can be identified.

Sah argues that disclosure reduces advisee bias when, as in medicine, there are [strong ethical norms](#) to "place patients first."⁴ Yet in some of Sah's

experimental designs in medical or financial contexts, medical advisors are warned that option A is clearly more beneficial to the audience than the advised option B, so the given medical advice (B) has no plausible deniability; but the financial advisors (who are in a role similar to advisors in CLM's experiments) have plausible deniability. The reader is left wondering: Is it the medical context or the lack of deniability that reduces bias? When Sah's designs correct this confound (by removing plausible deniability in both medical and business settings), the evidence merely suggests that disclosure reduces *intentional lying* when medical norms are manipulated. This is not evidence that disclosure reduces mere *bias* in medical settings. Since many medical contexts include problems of bias that go beyond intentional lying, the above flaws highlight the lack of external validity of Sah's prodisclosure experiments.

Real-World Problem of Conflicts of Interest

Most physicians (even biased ones) are not awake at night, thinking how to get rich by intentionally harming patients. Unfortunately, lay people's views of COIs (and even the view sometimes implied by prodisclosure research) often trade on a misconception that failure to properly navigate a COI is a problem of intentional corruption (ie, bad apples). This erroneous model depicts physicians as thinking, *I know that option A is best for my patients, but option B is best for my wallet. What should I do ... B?* This scenario gets the psychology wrong. COIs are not dangerous just for the intentionally corrupt Bernie Madoffs of the world. COIs are dangerous for people prone to unintentional bias—basically everyone.^{14,15}

The last 30 years of social science research has taught us that the human mind is simply not good at being objective.^{16,17} When physicians disclose a COI, it is not enough to trust that they *want* to be objective if they are psychologically incapable of being objective. Reducing bias is easier in black and white cases in which the physician knows for certain what is the best course for their patients. However, objectivity is more difficult in realistic gray areas in which the best course is uncertain. It's there that advisors might plausibly think that they can have their objective cake and eat it, too, thinking, *I know that option B is best for my wallet, but, of course, I put my patients first. The question is: What is best for my patients? That is less clear. It could also be ... B.*

References

1. Moore DA, Cain DM, Loewenstein G, Bazerman M, eds. *Conflicts of Interest: Challenges and Solutions in Business, Law, Medicine, and Public Policy*. New York, NY: Cambridge University Press; 2010.
2. Koch C, Schmidt C. Disclosing conflicts of interest—do experience and reputation matter? *Account Organ Soc*. 2010;35(1):95-107.
3. Church BK, Kuang X. Conflicts of interest, disclosure, and (costly) sanctions: experimental evidence. *J Legal Stud*. 2009;38(2):505-532.
4. Sah S. Conflict of interest disclosure as a reminder of professional norms: clients first! *Organ Behav Hum Decis Process*. 2019;154:62-79.
5. Cain DM, Loewenstein G, Moore DA. The dirt on coming clean: perverse effects of disclosing conflicts of interest. *J Legal Stud*. 2005;34(1):1-25.
6. Cain DM, Loewenstein G, Moore DA. When sunlight fails to disinfect: understanding the perverse effects of disclosing conflicts of interest. *J Consum Res*. 2011;37(5):836-857.
7. Sah S, Loewenstein G, Cain DM. The burden of disclosure: increased compliance with distrusted advice. *J Pers Soc Psychol*. 2013;104(2):289-304.
8. Sah S, Loewenstein G, Cain D. Insinuation anxiety: concern that advice rejection will signal distrust after conflict of interest disclosures. *Pers Soc Psychol Bull*. 2019;45(7):1099-1112.
9. Dana J, Cain DM, Dawes RM. What you don't know won't hurt me: costly (but quiet) exit in dictator games. *Organ Behav Hum Decis Process*. 2006;100(2):193-201.
10. Tversky A, Kahneman D. Judgment under uncertainty: heuristics and biases. *Science*. 1974;185(4157):1124-1131.
11. *PCAOB Public Hearing*, Houston, TX (October 18, 2012) (testimony of Robert Prentice, professor, University of Texas at Austin). https://pcaobus.org/Rulemaking/Docket037/ps_Prentice.pdf. Accessed December 30, 2019.
12. Bazerman MH, Moore DA. *Judgment in Managerial Decision Making*. 8th ed. New York, NY: Wiley; 2012.
13. Cain DM, Dana J, Newman GE. Giving versus giving in. *Acad Manag Ann*. 2014;8(1):505-533.
14. Cain DM, Detsky AS. Everyone's a little bit biased (even physicians). *JAMA*. 2008;299(24):2893-2895.
15. Loewenstein G, Sah S, Cain DM. The unintended consequences of conflict of interest disclosure. *JAMA*. 2012;307(7):669-670.

16. Kunda Z. The case for motivated reasoning. *Psych Bull.* 1990;108(3):480-498.
17. Moore DA, Tetlock PE, Tanlu L, Bazerman MH. Conflicts of interest and the case of auditor independence: moral seduction and strategic issue cycling. *Acad Manage Rev.* 2006;31(1):10-29.

Daylian M. Cain, PhD is a faculty member at the Yale School of Management in New Haven, Connecticut. He studies judgment and decision making and is an expert on the psychology of conflicts of interest.

Mohin Banker is a PhD student in marketing at the Yale School of Management in New Haven, Connecticut.

Citation

AMA J Ethics. 2020;22(3):E232-238.

DOI

10.1001/amajethics.2020.232.

Conflict of Interest Disclosure

The author(s) had no conflicts of interest to disclose.

The viewpoints expressed in this article are those of the author(s) and do not necessarily reflect the views and policies of the AMA.