

Anchoring on valuations and perceived informativeness[☆]Konstantinos Ioannidis^{*}

University of Cambridge, United Kingdom

CREED, University of Amsterdam & Tinbergen Institute, The Netherlands

ARTICLE INFO

Dataset link: [Github](https://github.com)

JEL classification:

D01

D91

C91

Keywords:

Anchoring

Replication

Experiment

ABSTRACT

Anchoring is a cognitive bias whereby individuals' decisions are influenced by an uninformative number, the anchor. Anchoring bias for valuations of goods has important implications for consumer decisions, but its robustness has been questioned by recent studies. We investigate the effect of the perceived informativeness of the anchor on valuations of goods. In an online experiment, we vary the amount of information about the process by which the anchor was determined, and hypothesise that the more information provided, the less scope is left for the anchor to be perceived as non-random/informative, thus mitigating anchoring effects. Our results provide evidence that the perceived informativeness of the anchor does affect anchoring effects. Contrary to our prediction, we find stronger anchoring effects when more information is presented.

1. Introduction

Anchoring is defined as the influence of an irrelevant cue on a judgement task. It originated in the pioneering work of [Tversky and Kahneman \(1974\)](#). In their experiments, a wheel of fortune containing numbers between 0 and 100 was spun and the resulting number –the anchor– was presented to subjects. Subjects were tasked to estimate percentages such as the percentage of African countries in the United Nations. They were first asked if they thought the percentage was higher or lower than the anchor, and then they provided their best estimate of the true value. The results showed that the anchor had a significant effect on their estimates, despite being randomly drawn. The effect has been documented to be robust in knowledge questions, where a true answer does exist, as well as probability judgements.

Though probability judgements are a key component of expected utility, anchoring research also challenged a fundamental assumption of standard economics, namely the assumption that agents have well defined valuations. In an influential study, [Ariely et al. \(2003\)](#) found that random anchors (the last two digits of the Social Security Number of subjects) affected the willingness to pay -WTP- of subjects for a range of ordinary goods (e.g., wine, chocolate, books). Subjects were first asked if they would be willing to buy the goods for a price equal to the

anchor before providing their incentivised WTP for the good. Subjects with above-median anchors provided 57 percent to 107 percent higher WTP than subjects with below-median anchors. The implications of the study are crucial for (among other domains) welfare evaluations as it demonstrated that preferences are malleable; valuations are not derived from fundamental values, but constructed on the spot under the influence of a random piece of information.

Given the important implications, the robustness of the results of [Ariely et al. \(2003\)](#) has been evaluated by subsequent studies. The results of this literature are mixed. While some papers documented similar effects ([Yoon & Fong, 2019](#); [Yoon et al., 2019](#)), a range of replication experiments found much smaller effects ([Bergman et al., 2010](#); [Sugden et al., 2013](#)), or even null effects ([Fudenberg et al., 2012](#); [Ioannidis et al., 2020](#); [Maniadi et al., 2014](#)). In order to better understand the puzzling pattern of the replication results, two recent meta-analyses aggregated existing evidence across experimental designs and investigated the circumstances under which anchoring effects are more likely to emerge.

[Li et al. \(2021\)](#) analysed 24 studies providing 53 experimental treatments. They found mixed evidence on whether anchoring effects are smaller for selling tasks, where willingness to accept -WTA- is elicited, compared to buying tasks,¹ larger anchoring effects in field and

[☆] Financial support from the Research Priority Area Behavioral Economics of the University of Amsterdam is gratefully acknowledged. The editor Levent Neyse and two anonymous referees provided constructive and highly valuable feedback during the review process. This paper was conditionally accepted as a registered report. The accepted design and preanalysis plan, can be found on [OSF](https://osf.io). All data and code can be found on [Github](https://github.com). The author would like to thank Tina Marjanov for her help in programming the experiment.

^{*} Correspondence to: University of Cambridge, United Kingdom.

E-mail address: ioannidis.a.konstantinos@gmail.com.

¹ The difference between WTP and WTA tasks is significant in some specifications and insignificant in other specifications.

classroom experiments compared to lab experiments, and no difference between incentivised and hypothetical elicitations of valuations. The key result which motivates the current study is related to the perceived informativeness of the anchor. In experiments where subjects were presented with an anchor without any explanation about how the anchor was determined, anchoring effects were larger compared to experiments where subjects were explicitly informed of the distribution of anchor values.

Ioannidis et al. (2020) analysed 19 studies providing 58 experimental treatments. They found that anchoring effects were larger for unfamiliar goods compared to familiar ordinary goods, and no difference between experiments which elicited WTA and experiments which elicited WTP. Most relevant for the current study, they analysed anchoring effects across different levels of informativeness of the anchor. They classified studies according to how informative the anchors can be perceived by the subjects. They used three categories: (i) informative anchors, where no information about the origin of the anchor was provided, (ii) semi-informative² anchors, where subjects were told the anchor is random, but no information about the distribution of the anchor was provided, and (iii) uninformative anchors, where subjects knew the anchor distribution and drew the anchor themselves. They found that anchoring effects were larger when the anchor is informative. Studies with transparently uninformative anchors are very few. According to Ioannidis et al. (2020) “although these studies are based on relatively many data, there are only few of them, and clearly more studies in this category are welcome”. Even if more studies were available, it is hard to compare results across different experimental designs as they may differ in more aspects than the ones included in a meta-analysis (e.g., different subject pool).

Closer to our research is Sugden et al. (2013) who also studied anchoring effects across a range of different anchor features such as the plausibility of the anchor, whether it was presented to subjects for a similar or for a dissimilar good, and whether subjects needed to actively search for the anchor.³ Among other results, they found evidence that active engagement in determining the anchor value increases anchoring effects. We cannot immediately conclude whether anchoring was observed due to the engagement itself or because the subjects believed the anchor value they had to discover was deliberately chosen by the experimenter, making it informative. Thus, their study is complementary to ours as it studies different ways of presenting the anchor whereas our study explicitly varies the information about the distribution of the anchor, and provides a cleaner setup to study how the perceived informativeness of the anchor moderates anchoring effects.

In our study, we vary the amount of information about the anchor distribution provided to our subjects across four experimental treatments. Our baseline treatment follows the canonical paradigm by presenting the anchor without any information about how it was determined. Subsequent treatments progressively provide more information that the anchor was drawn randomly, that the anchor was drawn randomly and the range of possible values, and that the anchor was drawn randomly together with the full distribution of possible values. We conjecture that the more information is provided about the randomness of the anchor (increasing across treatments), the less scope there is to perceive the anchor as informative (decreasing across treatments). Our main results provide evidence that the information about the anchor influences anchoring effects. However, contrary to our prediction, we find stronger anchoring effects when more information about the randomness of the anchor is provided. While acknowledging the difficulty of understanding our intriguing results, we discuss

conversational norms and the description-experience gap as plausible explanations.

We contribute to the anchoring literature in three ways. First, we provide direct experimental evidence to causally investigate the role of the perceived informativeness of the anchor on anchoring effects for valuations of goods. Second, as explained in the next section, all our treatments have been used in isolation in other studies. Consequently, our experiment provides direct replication for various designs by using a broader sample (online participants instead of students). Third, we provide additional data on the relatively understudied class of transparently uninformative anchor designs.

The rest of the paper continues as follows. In Section 2 we present our experimental design together with our predictions and a power analysis we run before the data collection. In Section 3 we present all our results, both between and within treatments. Finally, in Section 4 we position our results in the anchoring literature and discuss what we can learn from our study.

2. Experimental design

2.1. General design and implementation

The experiment consists of two tasks; the anchor task and the valuation task. Both tasks involve individual decisions about a lottery. The lottery is adapted from Sugden et al. (2013) with amounts approximately divided by a factor of five and the highest amount reduced to £2.88. Hence, subjects face a lottery with outcomes £2.88, £1.23 and £0.15 with probabilities 0.3, 0.5 and 0.2, respectively; its' expected value is £1.50.

The probabilities of the lottery are not explicitly showed to the subjects in numerical form. Instead, we show them on screen an urn with 100 coloured balls, whose colours correspond to the probabilities, and inform them of the amount they would get for each possible colour. The balls within the urn are sorted by colour to make the probabilities easy to visualise.⁴ The non-round outcomes and the fact that the exact probabilities are not shown numerically aims at making the calculation of the expected value cognitively demanding for subjects; thus minimising the chance of the expected value serving as an additional anchor.

In the anchor task, the subjects are asked if they are willing to sell the lottery for a price equal to the anchor. In all treatments, the anchor is a value between £0.00 and £3.00, drawn with uniform probability. The anchoring question is not incentivised.⁵ In the valuation task, subjects are asked to report the minimum price at which they would be willing to sell the lottery. The elicitation is incentivised via a price list. Subjects select whether they prefer to keep the lottery or sell the lottery for all prices between £0.00 and £3.00 in £0.10 increments. One of these decisions is randomly implemented for payment. This procedure is normatively equivalent to the Becker–DeGroot–Marschak mechanism (Becker et al., 1964). Overall, subjects receive a participation fee of £1.00 and additionally, depending on the randomly selected question from the price list, either the price for selling the lottery or the outcome of the lottery itself.

We use the same question for the anchor and the valuation task as this is the way it is typically operationalised in the anchoring literature (Ariely et al., 2003; Fudenberg et al., 2012; Maniadiis et al., 2014; Sugden et al., 2013). One interpretation of anchoring effects observed with this experimental design could be that anchoring is partially the result of preferences for consistency (Eyster, 2002; Falk &

² In Ioannidis et al. (2020) semi-informative anchors were called questionable anchors.

³ In the search treatments, subjects were asked to first find the lowest number in a matrix. The number became the anchor in the canonical comparative anchoring question.

⁴ The subjects can easily count the frequencies of each colour and use them to compute the expected value.

⁵ Anchoring manipulations in everyday life are typically not incentivised; they only serve to induce framing effects. Thus, keeping the anchor question hypothetical is externally valid (see also Sugden et al. (2013, Section 4.2)).

Zimmermann, 2017; Yariv, 2002). If subjects state their willingness to sell the lottery for a price of $\pounds X$, then preferences for consistency imply that their minimum WTA for the same lottery should not be higher than $\pounds X$.

Our study is run online, so it is not feasible to physically sell ordinary goods to the subjects while still maintaining their anonymity as this would require us to record their postal addresses. Using a lottery avoids this issue while remaining a valid design choice for two reasons. First, lotteries are cleaner to interpret than ordinary goods (like a bottle of wine or a keyboard used in previous studies) as they are truly private value commodities, and they do not have retail prices that subjects can refer to as homemade anchors. Second, Sugden et al. (2013) found similar effect sizes for lotteries and for ordinary goods, which reassures us that using lotteries should not affect our treatment comparisons.

We elicit WTA for three reasons. First, as found in the meta-analyses, there is little evidence that eliciting WTP or WTA affects anchoring effects. Second, eliciting WTA guarantees that subjects will receive a non-negative payment from either selling the lottery or from the lottery itself. On the contrary, eliciting WTP implies the amount paid to play the lottery is coming from their participation fee, which serves as initial endowment, and it could be perceived by the subjects as an additional anchor; thus confounding our anchoring manipulation. Third, though WTP is arguably more familiar to people, the popularity of WTA has increased with the rise of online marketplaces such as eBay or Facebook Marketplace.

2.2. Treatments

The experimental treatments vary how informative the anchor can be perceived. We do so by varying the information provided about the process by which the anchor was determined. They are presented here in decreasing order of perceived randomness. The four treatments are called NoInfo, RandomInfo, RangeInfo, and FullInfo.

In the first treatment (NoInfo), the anchors are presented to the subjects without any information about the origin of the anchor. Hence, the subjects may plausibly believe that the anchor is chosen non-randomly by the experimenter, and so is informative about the value of the lottery. Similar treatments are included in Ariely et al. (2003) and in Sugden et al. (2013).

The second treatment (RandomInfo) is a minimal deviation from the previous case. The subjects are again presented with an anchor without any information about the anchor distribution. The key difference is that they are explicitly informed that the anchor was drawn randomly. Similar treatments are included in Yoon et al. (2019) and in Yoon and Fong (2019). In those studies, the subjects were told they would draw a card with a number hidden behind, and this number was used as the anchor later. Despite knowing that the exact anchor was random, subjects could plausibly believe that the anchor is somehow correlated to a value, since they were unaware about the distribution of anchor values.

The third treatment (RangeInfo) goes one step further by informing the subjects about the range, but not the distribution of the possible anchor values. This treatment is inspired by evidence suggesting that the elicitation format, and more specifically the range of admitted prices, affects bidding (Spann et al., 2005). In the experiment of Ioannidis et al. (2020), the modal response for the valuation of a bottle of wine was $\pounds 5$, even though the retail price of the wines used in the experiment ranged between $\pounds 6.00$ and $\pounds 7.50$. It is plausible that the midpoint of the range of possible anchors ($\pounds 0.00$ to $\pounds 9.90$) was interpreted by subjects as an indication of the retail price. This treatment resembles the designs of Chapman and Johnson (1999) and Ariely et al. (2003), where the anchor was created from the last two digits of the social security numbers of the subjects. In a post-experiment questionnaire, one third of the subjects in Chapman and Johnson (1999) mentioned they thought the anchor was informative.

The fourth treatment (FullInfo) fully reveals the distribution of anchor values to subjects. Subjects in this treatment should not attach any information value to the presented anchor. This treatment is similar to the full information anchor procedure of Fudenberg et al. (2012), where the software Excel produced a random number when subjects clicked a button, and of Ioannidis et al. (2020), where subjects rolled a 10-sided die twice to create the anchor. To keep this treatment comparable with the other treatments, our subjects do not produce the anchor themselves, but they are fully informed about the anchor distribution.

The treatment variations are implemented by providing progressively more information about the origin of the anchor to subjects. In NoInfo, the subjects are asked the question “Would you be willing to sell the lottery you own for a price of $\pounds X$?”. In RandomInfo, we additionally inform them that X was randomly drawn before the experiment. In RangeInfo, we additionally specify that X was randomly drawn between $\pounds 0.00$ and $\pounds 3.00$ before the experiment. In FullInfo, we additionally mention that all values between $\pounds 0.00$ and $\pounds 3.00$ had the same probability of being drawn. On top of standard demographics such as age, and gender, we ask subjects whether they think X provided information about their valuation of the lottery, and whether they believe the experimenters wanted X to affect their valuation decision.

2.3. Predictions

Our primary prediction is that anchoring effects will differ between treatments. To test our prediction, we use a test for equality of correlations across all treatments (Jennrich, 1970). The null hypothesis is that all correlations between anchor and WTA are equal; the alternative is that correlations differ.⁶

Our secondary predictions focus on whether we observe significant anchoring effects within each treatment. Following popular approaches in the anchoring literature, we formally test for anchoring effects in three ways for each treatment. First, we perform a Mann–Whitney rank-sum test on WTA between high and low anchor groups. Second, we compute the Pearson correlation between anchor and WTA and test whether it is positive. Third, we estimate tobit regressions of WTA on anchor, while also controlling for subject demographics. Our prediction is that the likelihood of observing a significant effect will be monotonous across treatments. More specifically, we expect the likelihood to be decreasing when more information is provided.

Finally, given that our predictions are directed, we make pairwise comparisons of the magnitude of the effects between treatments. To do so, we end our analysis with two more tests using data from all our treatments. First, a tobit regression of WTA on the interaction of anchor and treatment. Second, we use difference-in-differences tests to compare WTA between low and high anchor groups and between treatments. We expect a larger anchoring effect (so significantly positive DID coefficient) when comparing a treatment with less information with a treatment with more information.

2.4. Power analysis

To calculate the necessary sample size to be sufficiently powered in our study, we use a Monte-Carlo simulation. The simulation repeats the following steps multiple times. First, we generate a dataset assuming the alternative hypothesis is true. We rely on the two meta-analyses to indicate the range of correlations under the alternative. Second, we test

⁶ The test does not require all pairwise comparisons to be significant, but rejects the null hypothesis if there is sufficient heterogeneity in the correlations; a feature shared with well known statistical tests like ANOVA, which also tests if all group means are equal or not without requiring that all means differ. The test is included in Stata under the command *mvtest correlations*.

the null hypothesis that all correlations are equal using the simulated dataset. Third, we record whether the null hypothesis is rejected or not. The mean rejection rate provides an estimate of the statistical power of our test. Based on our power analysis (see [Appendix A](#) for details) and taking into account that online experiments have higher levels of noise, we aimed at collecting at least 150 observations per treatment.

2.5. Demographics

The experiment was programmed in Qualtrics and run online using Prolific ([Palan & Schitter, 2018](#)) in February of 2023. In total we recruited 681 subjects from the UK; 151 in NoInfo, 191 in RandomInfo, 153 in RangeInfo, and 186 in FullInfo. Our subjects were balanced in terms of gender (54.8% males, 44.2% females, 1.0% other) and highest attained education level (42.9% high school, 57.1% university), and are on average 40.22 years old (sd = 13.98, min = 18, max = 79). The demographic characteristics were also balanced across treatments.

3. Results

We begin with our primary hypothesis of equality of correlations between WTA and anchor across all treatments. The raw correlations are -0.011 in NoInfo, 0.141 in RandomInfo, 0.180 in RangeInfo, and 0.288 in FullInfo. Equality of correlations is rejected by the data ($\chi^2 = 7.93, p = 0.047, N = 681$). It is important to notice that the correlations between RandomInfo and RangeInfo treatments are (expectedly) close.⁷ Thus, we additionally test for equality of correlations by either dropping RangeInfo ($\chi^2 = 7.78, p = 0.020, N = 528$), dropping RandomInfo ($\chi^2 = 7.90, p = 0.019, N = 490$), merging RandomInfo and RangeInfo ($\chi^2 = 7.81, p = 0.020, N = 681$), or randomly selecting half of the observations from RandomInfo and RangeInfo (averaged across 100,000 random subsets: $p = 0.035, N = 509$). All specifications provide support for our first result.

Result 1. *Anchoring effects are influenced by the information presented about the anchor.*

We proceed by testing our secondary predictions of anchoring effects within each treatment. For each treatment separately, two anchor groups are created, namely the low anchor group and the high anchor group. The groups are determined by a median split on the anchors.⁸ We visualise our results in [Fig. 1](#), which shows boxplots of WTA for low and high anchor groups for all treatments. We make two observations before we proceed with formal tests. First, we can already notice that anchoring is increasing as we move along the x-axis of the figure, since the gap in median WTA between low and high anchor groups is increasing. Second, we observe that low anchors affected WTA more strongly, especially in the FullInfo treatment, while WTA from high anchors do not differ much across treatments.

We now formally test for anchoring effects for each treatment. First, we use Mann–Whitney rank-sum tests. The prediction is that the WTA in the high anchor group will be higher than the WTA in the low anchor group. We find no anchoring effects in NoInfo ($z = 0.117, p = 0.907, N = 151$), significant anchoring effects in RandomInfo ($z = 2.269, p = 0.023, N = 191$), marginally significant anchoring effects in RangeInfo ($z = 1.894, p = 0.058, N = 153$), and significant anchoring effects in FullInfo ($z = 3.908, p = 0.001, N = 186$). Second, we test whether the correlations between WTA and anchor are significant. The correlation

⁷ Our informed prediction was that those two treatments would likely produce similar anchoring effects.

⁸ Assigning all subjects with anchor lower than 1.5 to the low anchor group and all subjects with anchor higher than 1.5 to the high anchor group will change the classification into anchor groups for only 20 out of 681 subjects. Thus, our results are robust to this alternative way of determining anchor groups.

Table 1

Pairwise comparisons of anchoring effects between treatments.

	NoInfo vs. RandomInfo	NoInfo vs. RangeInfo	NoInfo vs. FullInfo	RandomInfo vs. RangeInfo	RandomInfo vs. FullInfo	RangeInfo vs. FullInfo
Model 1	0.82 (0.443)	1.11 (0.330)	3.62** (0.028)	0.11 (0.897)	2.05 (0.129)	1.17 (0.309)
Model 2	1.33 (0.185)	1.21 (0.228)	2.51** (0.013)	0.05 (0.962)	1.22 (0.223)	1.23 (0.219)

Model 1: F-values from coefficient comparisons and p-values in parentheses.

Model 2: t-values from DID interaction term.

Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

is not significant in NoInfo ($r = -0.011, p = 0.896, N = 151$), marginally significant in RandomInfo ($r = 0.141, p = 0.052, N = 191$), significant in RangeInfo ($r = 0.180, p = 0.026, N = 153$), significant in FullInfo ($r = 0.288, p = 0.001, N = 186$). Third, we test for anchoring effects using separate tobit regressions of WTA on anchor for each treatment. We obtain the same pattern with insignificant anchoring effects in NoInfo ($b = 0.002, SE = 0.075, CI = [-0.146, 0.150], t = 0.02, p = 0.980, N = 151$), marginally significant in RandomInfo ($b = 0.118, SE = 0.064, CI = [-0.007, 0.244], t = 0.86, p = 0.065, N = 191$), significant in RangeInfo ($b = 0.155, SE = 0.073, CI = [0.011, 0.298], t = 2.13, p = 0.035, N = 153$), significant in FullInfo ($b = 0.251, SE = 0.064, CI = [0.126, 0.377], t = 3.95, p = 0.001, N = 186$).

Result 2. *There are no anchoring effects in NoInfo, some evidence of anchoring effects in RandomInfo and RangeInfo, and significant anchoring effects in FullInfo.*

We end our analysis with pairwise comparisons of anchoring effects between treatments. We first estimate a tobit regression of WTA on anchor, treatment, and their interaction effect. We then proceed by testing the joint hypothesis that the treatment effects and the interaction effects are equal for pairs of treatments.⁹ The results are shown in Model 1 in [Table 1](#). We also estimate a difference-in-differences model to compare WTA between low and high anchor groups and between treatments. Model 2 in the table shows the results of the DID estimations. The overall pattern provides evidence for significant differences in anchoring effects only between our two most extreme information treatments, namely NoInfo and FullInfo. However the difference is in the opposite direction of what we expected.

Result 3. *Anchoring effects differ between FullInfo and NoInfo. Anchoring effects do not differ between any other two treatments.*

We end our result section by reporting on the post-experiment survey questions. There is no difference between treatments on the self-reported beliefs about whether subjects were affected by the anchor (ranging between 60.9% and 70.4%), and whether subjects believed we wanted them to be affected (ranging between 78.1% and 81.1%). Those questions are not significant in any of our previous regressions and our results are robust to including them as controls.

4. Concluding discussion

Despite the vast literature on anchoring, the phenomenon remains debated and impartially understood. Our study aims to investigate how the perceived randomness of an anchor affects how informative the anchor is perceived to be, and consequently how strong anchoring effects are. We provide evidence that the perceived informativeness does affect anchoring effects. Thus, our results help our understanding

⁹ The results are qualitatively similar if we make pairwise comparisons between main treatment effects only or interaction effects only.

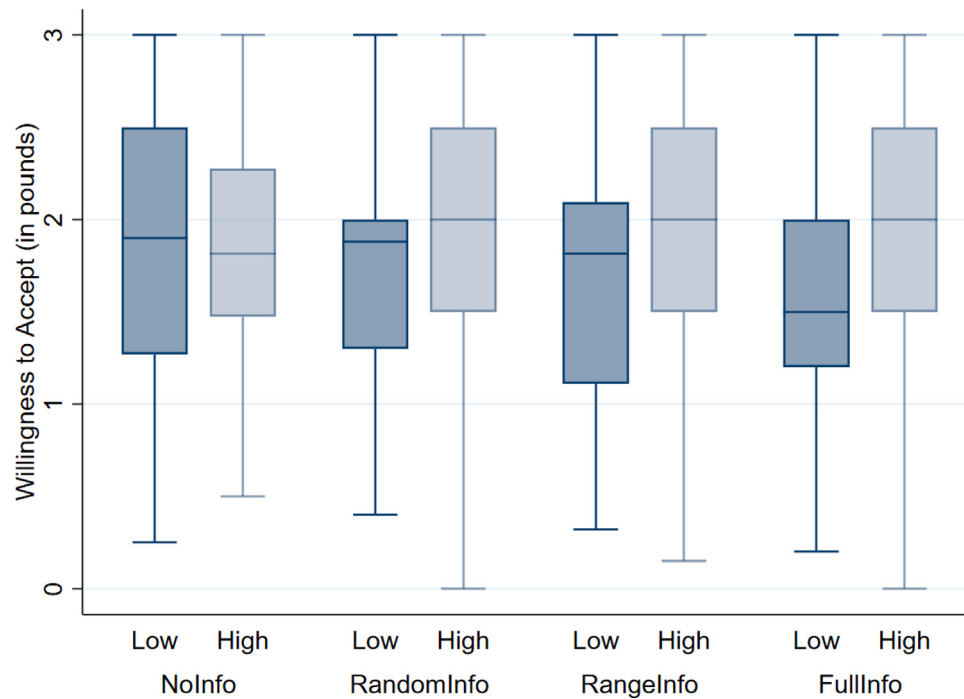


Fig. 1. Boxplots of WTA between low and high anchor groups and over treatments.

of the circumstances under which anchoring is more likely to emerge, while remaining inconclusive on why it is so.

Our study differs from the majority of previous studies in two aspects. First, it is run online which often results in higher noise and weaker incentives. Although we are sufficiently powered to detect effect sizes of a magnitude similar to lab studies, we cannot rule out that running out study online may have affected our results. However, we believe this is a strength of our study as one of our goals was to replicate previous designs in a broader sample. Second, we provide a visualisation of the lottery (see [Appendix B.2](#)) which can facilitate a better understanding of probabilities, but does differ from other studies that provided a written description of the probabilities.

Our study was designed to test how perceived informativeness affects anchoring, and so we cannot be conclusive about other mechanisms. However, in the remaining of the section, we focus on two of our results which were surprising, and we discuss plausible explanations which are consistent with our findings.

The direction of anchoring effect in our study is unexpected. While we expected anchoring to be mitigated when more information about the anchor is provided, we find the opposite. Our subjects react to the amount rather than the content of information, which implies that perceived informativeness and perceived randomness of the anchor are not interchangeable. We note that our results are not due to insufficient attention of subjects as we observe a clear pattern. The theory of conversational norms ([Grice, 1975](#)) can help us navigate our puzzling results. Conversational norms, also known as the cooperative principle, is a linguistic theory which suggests that people rely on four maxims in order to achieve effective communication: informativeness, truthfulness, relevance, and clarity. Thus, in every day conversation, information is never provided with the intention to be ignored; senders of information aim to influence the information set and the beliefs of receivers. As we know from the cheap talk literature ([Crawford & Sobel, 1982](#)), it may be optimal for a receiver to ignore the information from a sender if their preferences are strongly misaligned, but the goal of the sender is not to be ignored. For anchoring, the theory implies

that subjects reasoned that the experimenter would not have presented the anchors if they were not informative; a conjecture also discussed in [Chapman and Johnson \(1999\)](#) and in [Sugden et al. \(2013\)](#). In our setting, conversational norms would imply that the more information we provide about the anchor, the harder it is for subjects receiving this information to ignore it. While our results are consistent with this theory, future experiments varying orthogonally the amount and the content of information about the anchor are needed. If conversational norms are indeed at play here, we would expect anchors to be perceived as more informative when more information is provided, *irrespective* of whether the provided information is related to the distribution of the anchor or not.

Looking at our treatments in isolation, the most surprising result is the strong anchoring effect we document when all information about the anchor is presented. When designing the experiment, we believed our FullInfo treatment to be isomorphic to previous experiments where the full distribution of the anchor was known to subjects. While those studies found no anchoring effects, we find strong anchoring effects. We believe the contradicting results could be due to the description-experience gap ([Hertwig et al., 2004](#); [Hertwig & Erev, 2009](#)). The description-experience gap suggests that behaviour may differ depending on whether the probabilities of uncertain outcomes are described or experienced. According to this theory, subjects may perceive the randomness of the anchor differently if they *experience* it themselves (e.g., when rolling a die to generate the anchor as in [Ioannidis et al. \(2020\)](#)) compared to our current setting where the full distribution of the possible anchor values is *described* to the subjects.

Data availability

All data and code can be found on [Github](#).

Appendix A. Power analysis for equality of correlations test

Relying on the meta-analyses, we estimate the correlations of our treatments assuming treatment effects are real. For FullInfo treatment,

Table 2

Power calculations.

(a) Significance level $\alpha = 0.05$

Power	Sample	SD	Predicted correlations			
			FullInfo	RangeInfo	RandomInfo	NoInfo
0.763	400	0.25	0.03	0.16	0.20	0.32
0.827	450	0.25	0.03	0.16	0.20	0.32
0.855	500	0.25	0.03	0.16	0.20	0.32
0.901	550	0.25	0.03	0.16	0.20	0.32
0.926	600	0.25	0.03	0.16	0.20	0.32

(b) Significance level $\alpha = 0.10$

Power	Sample	SD	Predicted correlations			
			FullInfo	RangeInfo	RandomInfo	NoInfo
0.852	400	0.25	0.03	0.16	0.20	0.32
0.893	450	0.25	0.03	0.16	0.20	0.32
0.911	500	0.25	0.03	0.16	0.20	0.32
0.956	550	0.25	0.03	0.16	0.20	0.32
0.962	600	0.25	0.03	0.16	0.20	0.32

Notes: SD = Standard deviation of main outcome.

estimates range between -0.01 and 0.06 . For RandomInfo and RangeInfo treatments, we have no clear indication as the studies in the meta-analyses do not distinguish between those two categories. Assuming a small difference between them, we use estimates ranging between 0.10 – 0.26 and 0.13 – 0.27 respectively. For NoInfo, estimates range between 0.22 and 0.43 . We use the midpoint of the estimates to create the simulated datasets.

Table 2 provides power calculations for different significance levels and sample sizes. For significance level of 0.05 , we need at least 450 subjects to achieve a power of 0.80 , whereas with a significance level of 0.10 we can already exceed power of 0.80 with less than 400 subjects.¹⁰

Appendix B. Instructions

B.1. Welcome screen

You are participating in a study of the University of Amsterdam. Your participation in this study is confidential and your identity will not be stored with your data.

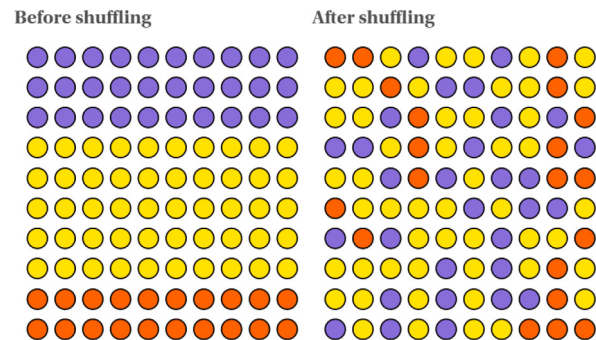
You will not be asked to provide any personally identifying information. Your answers will be used only for research purposes and in ways that will not reveal who you are. The results will be reported in aggregate form only, and cannot be identified individually. Any information that could identify you will be removed before data is shared with other researchers or results are made public. By participating in this study, you consent to the data being used for this purpose.

Your participation in this research is entirely voluntary and you have the right to withdraw consent at any time by closing your browser. The goal of this study is to assess how people make decisions. You receive a guaranteed payment of $\text{€}1.00$ for participating in the study. Further, you can earn substantial money as a bonus. Your bonus payment will depend on your decisions. The study is expected to take 10 min.

¹⁰ We repeated the calculations by using the lower and upper bounds of the correlation ranges in all possible combinations (16 per treatment) as well as by varying the standard deviation of the main outcome. Power exceeded 0.80 for the vast majority of combinations.

B.2. Lottery description

The study involves making a decision about a lottery with three possible outcomes. We provide you with a visualisation of the lottery below. Blue balls are worth $\text{€}2.88$, yellow balls are worth $\text{€}1.23$, and red balls are worth $\text{€}0.15$. After shuffling the balls, the computer will randomly draw a ball and this will be the outcome of the lottery.



In the remainder of the study, consider this lottery yours. Your task is to decide between two options.

- Option 1: play the lottery If you play the lottery, the computer will draw a ball from the urn. Your bonus payment will be the amount corresponding to the ball that was drawn.
- Option 2: sell the lottery If you sell the lottery, your bonus payment will be the amount you sold it for and the lottery will not be implemented.

B.3. Anchoring question

The answer to the following question will not affect your bonus payment.¹¹ Your bonus payment will be determined by your decision in the next screen.

- The price of the lottery is $\text{€}X$. [NoInfo, RandomInfo, RangeInfo, FullInfo]
- The price of $\text{€}X$ was randomly drawn before the experiment. [RandomInfo, RangeInfo, FullInfo]
- The possible values ranged between $\text{€}0.00$ and $\text{€}3.00$. [RangeInfo, FullInfo]
- All values had the same probability of being drawn. [FullInfo]

Please indicate what you prefer:

For a price of $\text{€}X$: Play the Lottery OR Sell the lottery.

B.4. WTA elicitation

Your bonus payment will be determined by your decision on this task.

Your task is to indicate whether you would prefer to keep the lottery or to sell the lottery for each price below. One of those prices will randomly be drawn and your decision on whether to keep or sell the lottery for that price will determine your bonus payment.

We remind you that:

- if you keep the lottery, the computer will implement it and your bonus payment will be the amount corresponding to the lottery outcome.
- if you sell the lottery, your bonus payment will be the price you sold it for and the lottery will not be implemented.

¹¹ In square brackets we indicate in which treatment each of the bulletpoints was actually shown to our subjects.

To make the decision easier for you, we provide you with the slider below. Once you decide the minimum price at which you are willing to sell the lottery, please move the slider to that price. When you do that, the system will auto-complete all the choices for you so that you keep the lottery for all prices lower than your selected price, and that you sell the lottery for all prices higher than your selected price.

Please select the minimum price at which you are willing to sell the lottery.

Slider was provided here

For a price of £0.00: Play the Lottery OR Sell the lottery.

For a price of £0.10: Play the Lottery OR Sell the lottery.

For a price of £0.20: Play the Lottery OR Sell the lottery.

For a price of £0.30: Play the Lottery OR Sell the lottery.

For a price of £0.40: Play the Lottery OR Sell the lottery.

For a price of £0.50: Play the Lottery OR Sell the lottery.

For a price of £0.60: Play the Lottery OR Sell the lottery.

For a price of £0.70: Play the Lottery OR Sell the lottery.

For a price of £0.80: Play the Lottery OR Sell the lottery.

For a price of £0.90: Play the Lottery OR Sell the lottery.

For a price of £1.00: Play the Lottery OR Sell the lottery.

For a price of £1.10: Play the Lottery OR Sell the lottery.

For a price of £1.20: Play the Lottery OR Sell the lottery.

For a price of £1.30: Play the Lottery OR Sell the lottery.

For a price of £1.40: Play the Lottery OR Sell the lottery.

For a price of £1.50: Play the Lottery OR Sell the lottery.

For a price of £1.60: Play the Lottery OR Sell the lottery.

For a price of £1.70: Play the Lottery OR Sell the lottery.

For a price of £1.80: Play the Lottery OR Sell the lottery.

For a price of £1.90: Play the Lottery OR Sell the lottery.

For a price of £2.00: Play the Lottery OR Sell the lottery.

For a price of £2.10: Play the Lottery OR Sell the lottery.

For a price of £2.20: Play the Lottery OR Sell the lottery.

For a price of £2.30: Play the Lottery OR Sell the lottery.

For a price of £2.40: Play the Lottery OR Sell the lottery.

For a price of £2.50: Play the Lottery OR Sell the lottery.

For a price of £2.60: Play the Lottery OR Sell the lottery.

For a price of £2.70: Play the Lottery OR Sell the lottery.

For a price of £2.80: Play the Lottery OR Sell the lottery.

For a price of £2.90: Play the Lottery OR Sell the lottery.

For a price of £3.00: Play the Lottery OR Sell the lottery.

B.5. Post-experiment survey

Earlier in the study you were asked if you would hypothetically be willing to sell the lottery for a price of £X. Was £X informative when you were choosing the minimum price at which you are willing to sell the lottery (the slider task)? [Yes] [No]

Do you think the experimenters wanted £X to influence your choice of the minimum price at which you are willing to sell the lottery (the slider task)? [Yes] [No]

Please indicate your age.

Please indicate your gender.

Please indicate your education level.

B.6. (Example of) Feedback screen

This page will explain to you how your final payment is calculated. You will receive £1.00 for participating in the study. The randomly drawn price was £Y. For this price, you indicated you were willing to sell the lottery. Hence, your bonus payment is equal to £Y. Your total payment is £1 + Y.

Thank you for participating in this study.

References

- Ariely, D., Loewenstein, G., & Prelec, D. (2003). Coherent arbitrariness: Stable demand curves without stable preferences. *Quarterly Journal of Economics*, 118(1), 73–106.
- Becker, G. M., DeGroot, M. H., & Marschak, J. (1964). Measuring utility by a single-response sequential method. *Systems Research and Behavioral Science*, 9(3), 226–232.
- Bergman, O., Ellingsen, T., Johannesson, M., & Svensson, C. (2010). Anchoring and cognitive ability. *Economics Letters*, 107(1), 66–68.
- Chapman, G. B., & Johnson, E. J. (1999). Anchoring, activation, and the construction of values. *Organizational Behavior and Human Decision Processes*, 79(2), 115–153.
- Crawford, V. P., & Sobel, J. (1982). Strategic information transmission. *Econometrica*, 50(6), 1431–1451.
- Eyster, E. (2002). *Rationalizing the past: a taste for consistency*. Nuffield College Mimeograph.
- Falk, A., & Zimmermann, F. (2017). Consistency as a signal of skills. *Management Science*, 63(7), 2197–2210.
- Fudenberg, D., Levine, D. K., & Maniadis, Z. (2012). On the robustness of anchoring effects in WTP and WTA experiments. *American Economic Journal: Microeconomics*, 4(2), 131–145.
- Grice, H. P. (1975). Logic and conversation. In *Speech acts* (pp. 41–58). Brill.
- Hertwig, R., Barron, G., Weber, E. U., & Erev, I. (2004). Decisions from experience and the effect of rare events in risky choice. *Psychological Science*, 15(8), 534–539.
- Hertwig, R., & Erev, I. (2009). The description–experience gap in risky choice. *Trends in Cognitive Sciences*, 13(12), 517–523.
- Ioannidis, K., Offerman, T., & Sloof, R. (2020). On the effect of anchoring on valuations when the anchor is transparently uninformative. *Journal of the Economic Science Association*, 6(1), 77–94.
- Jennrich, R. I. (1970). An asymptotic χ^2 test for the equality of two correlation matrices. *Journal of the American Statistical Association*, 65(330), 904–912.
- Li, L., Maniadis, Z., & Sedikides, C. (2021). Anchoring in economics: a meta-analysis of studies on willingness-to-pay and willingness-to-accept. *Journal of Behavioral and Experimental Economics*, 90, Article 101629.
- Maniadis, Z., Tufano, F., & List, J. A. (2014). One swallow doesn't make a summer: New evidence on anchoring effects. *American Economic Review*, 104(1), 277–290.
- Palan, S., & Schitter, C. (2018). Prolific. ac—A subject pool for online experiments. *Journal of Behavioral and Experimental Finance*, 17, 22–27.
- Spann, M., Bernhardt, M., Haubl, G., & Skiera, B. (2005). It's all in how you ask: Effects of bid-elicitation format on bidding behavior in reverse-pricing markets. In *Proceedings of the 26th annual conference of the society for judgment and decision making*. SJDM, SJDM.
- Sugden, R., Zheng, J., & Zizzo, D. J. (2013). Not all anchors are created equal. *Journal of Economic Psychology*, 39(1), 21–31.
- Tversky, A., & Kahneman, D. (1974). Judgment under uncertainty: Heuristics and biases. *Science*, 185(4157), 1124–1131.
- Yariv, L. (2002). *I'll see it when I believe it? A simple model of cognitive consistency: Cowles foundation discussion papers*. 1616.
- Yoon, S., & Fong, N. (2019). Uninformative anchors have persistent effects on valuation judgments. *Journal of Consumer Psychology*, 29(3), 391–410.
- Yoon, S., Fong, N. M., & Dimoka, A. (2019). The robustness of anchoring effects on preferential judgments. *Judgment and Decision Making*, 14(4), 470–487.