

Online Appendix for: What to blame? Self-serving attribution bias with multi-dimensional uncertainty

Alexander Coutts

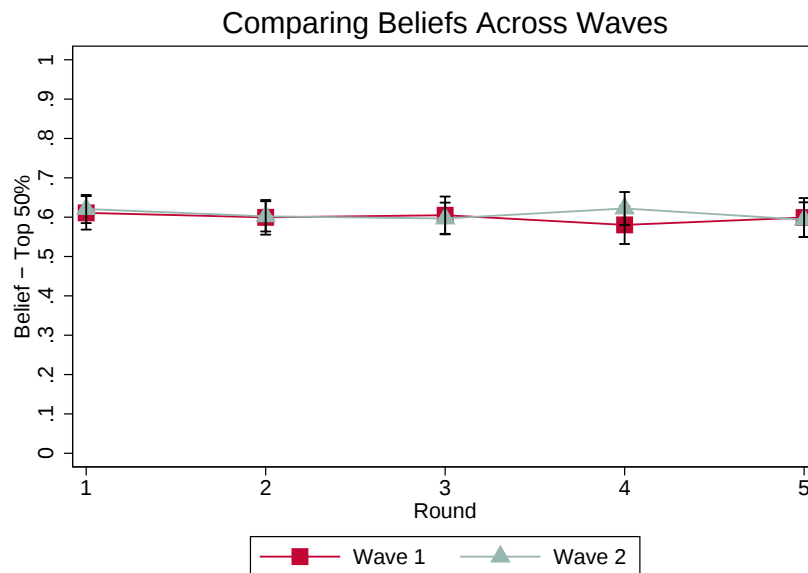
Leonie Gerhards

Zahra Murad

1 Beliefs by Wave

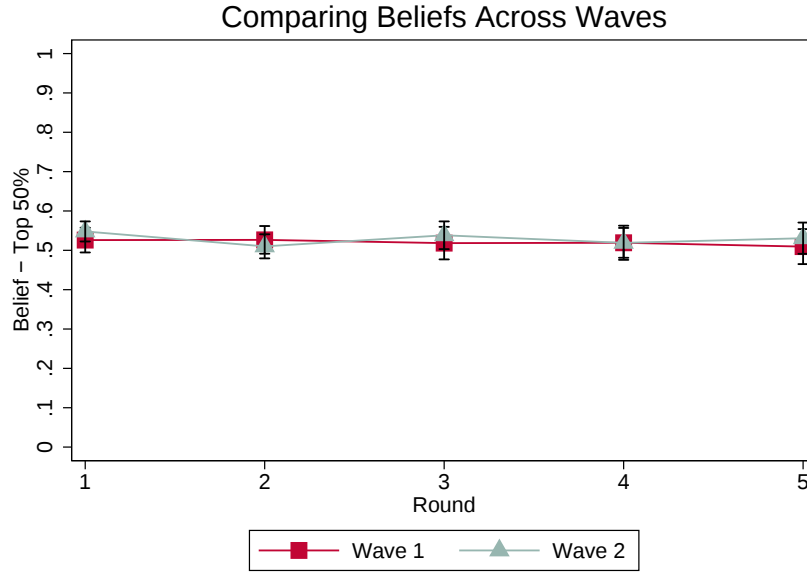
Figures OA.1 and OA.2 present average beliefs about teammate 1 and 2 being in the top half of performers respectively. Due to slight differences in the belief elicitation we may be concerned that these resulted in differences in reported beliefs. From these figures we can see that there are no differences in beliefs across waves, neither for teammate 1 nor 2.

Figure OA.1: Beliefs about teammate 1 across waves 1 and 2



Average beliefs about teammate 1 for rounds 1 to 5, split by wave. 95% confidence intervals shown.

Figure OA.2: Beliefs about teammate 2 across waves 1 and 2



Average beliefs about teammate 2 for rounds 1 to 5, split by wave. 95% confidence intervals shown.

2 Hard Easy Effect

In Part 1 of the experiment, subjects took either an easy or a hard version of the test. The questions had very similar scope, but differed in difficulty, which was assessed both by the authors and confirmed by piloting of the test questions. With respect to relative comparisons, the hard-easy effect, see [Larrick et al. \(2007\)](#) and [Moore and Small \(2007\)](#), predicts that individuals will predict that their performance relative to others is greater on easy tasks relative to hard tasks.¹

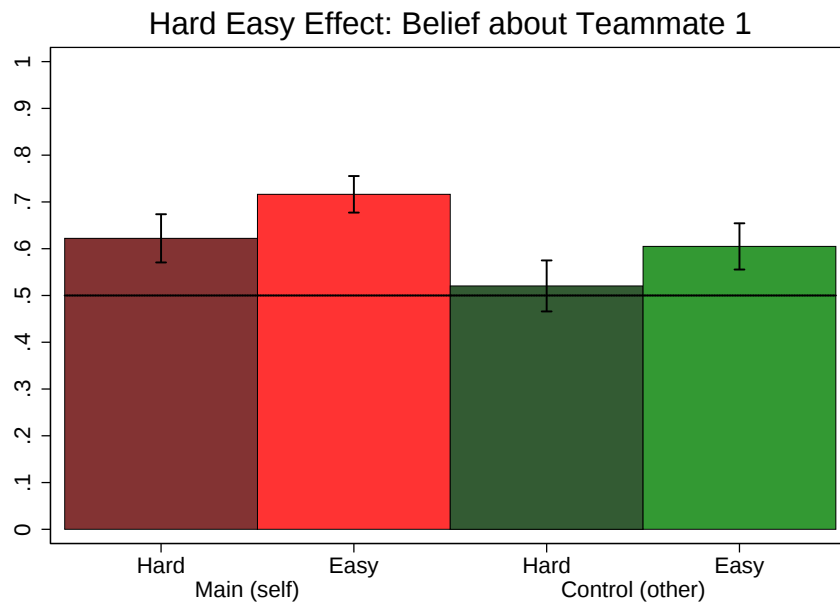
Figure OA.3 presents the prior beliefs, before receiving any feedback, about the probability that teammate 1 was in the top half of performers. Recall that in Main treatment, teammate 1 was the subject themselves, while in Control it was a third party (whose responses were visible to the subject). The hard-easy manipulation is successful, as individuals believe they are in the top half with 72% probability when the test was easy, compared with 62% when the test was hard (Wilcoxon rank-sum p-value: 0.037). Interestingly, the hard-easy effect carries through to the Control treatment, where individuals are assessing the performance of a stranger in the position of teammate 1. There initial beliefs for the hard test are 61% and for the easy test are 52% (Wilcoxon rank-sum p-value:

¹Overconfidence in relative performance is often referred to as over-placement. For estimation about absolute performance, this hard-easy effect is reversed, i.e. there is more over-estimation in hard tasks.

0.047).²

While we found hard-easy effects for teammate 1, regardless of whether the subject was in the team, Figure OA.4 shows quite clearly that there are no effects for teammate 2. This is true in Main, when the subject is herself teammate 1, and also for Control, when the subject is not a member of the team. Regarding updating beliefs (not shown) no clear patterns across hard and easy tests emerge for both teammates. For teammate 1 there is positive asymmetry of similar magnitude for both test versions (at the 5% and 10% level, for hard and easy respectively), while for teammate 2 there is positive asymmetry for the hard test (significant at the 10% level), but no evidence of positive asymmetry for the easy test.

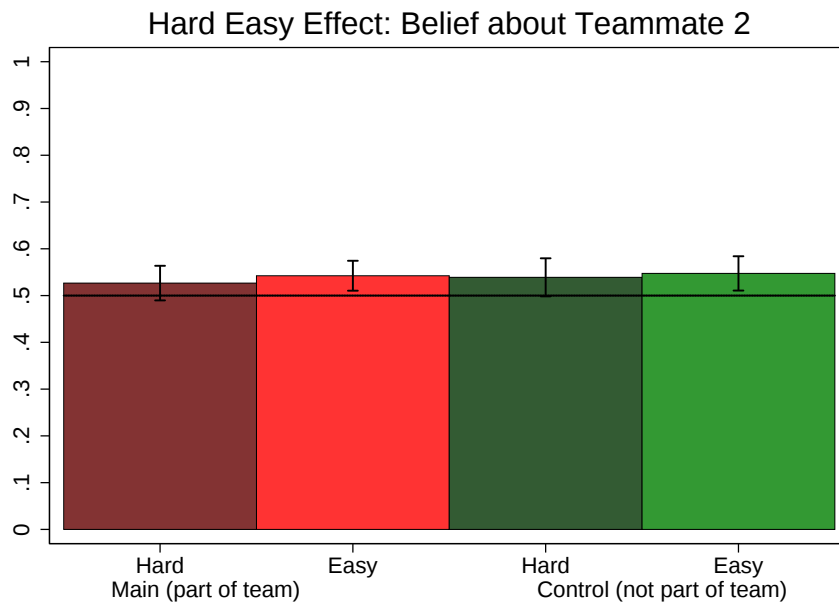
Figure OA.3



Examining the hard-easy effect in initial beliefs about teammate 1, before receiving feedback. In Main treatment subject is teammate 1. In Control, subject is not on the team, but is aware of the responses on the test of teammate 1.

²Additionally, we continue to find evidence of overconfidence comparing Main and Control within either hard or easy sub-samples, Wilcoxon rank-sum p-values are 0.014 and 0.006 respectively.

Figure OA.4



Examining the hard-easy effect in initial beliefs about teammate 2, before receiving feedback. In Main treatment subject is on the team, as teammate 1. In Control, subject is not on the team. In both cases subject is aware of only the number of questions attempted by teammate 2.

3 Gender Differences

3.1 Prior Beliefs

We first note that our sample is balanced across gender, with 52% women. Figure OA.5 (a) presents prior beliefs of males and female subjects respectively for prior beliefs about the probability of teammate 1 being in the top half, for Main treatment (their own performance) and for the Control treatment (performance of other). It appears that men are dramatically more overconfident than women about their own performance (Main), both relative to the Bayesian prediction of 0.5 and the to the Control treatment. However, there are also performance differences on the test, at lower test score levels. Men scored significantly better than women overall, though this difference disappears if we eliminate the bottom one-third of performers.³

To account for these potential differences we further split the sample into teammate 1's

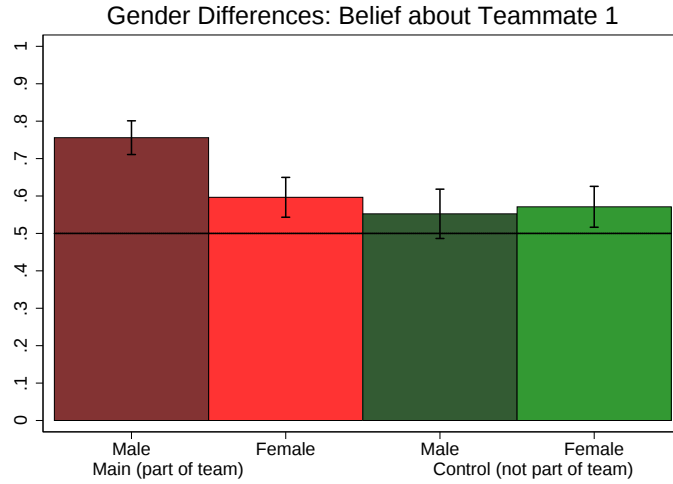
³We note that these were not true IQ test but ad-hoc tests created with various logic and trivia questions. In our framing however, we emphasized that, "Questions similar to these are often used to measure a person's general intelligence (IQ)". It is possible that this generated pressure which may have disproportionately affected females rather than males, i.e. stereotype threat, see [Spencer et al. \(2016\)](#). Consistent with such an explanation is that there are no gender differences for the top two-thirds of performers.

who scored in the top half (those ranked 1-10), and those who scored in the bottom half (those ranked 11-20). Results should thus be interpreted cautiously, due to the smaller samples. Figures OA.5 (b) and (c) present prior beliefs, by gender, for these respective groups. Because we are conditioning on actual performance of teammate 1, the Bayesian prediction is in fact 1 for those in the top half, and 0 for those in the bottom half.⁴ Thus these results show the distribution of underconfidence and overconfidence respectively. For those in the top, Figure OA.5 (b) shows significant, though more moderate underconfidence. Males are more confident than females both regarding own performance (Main, Wilcoxon rank-sum test p-value: 0.004) and other performance (Control, Wilcoxon rank-sum test p-value 0.016). Interestingly, neither males nor females in Main are more confident about their own performance, compared to the same respective genders in Control (regarding other performance), thus this underconfidence likely reflects only the lack of full information. For those in the bottom Figure OA.5 (c) shows that there is significant overconfidence. Males *in the bottom half* believe on average there is a 67% probability they are in the top half. For females this is still strikingly large, at 54%. The gender difference is significant at the 5% level, Wilcoxon rank-sum test p-value 0.030. While this already indicates overconfidence, we can also note that the estimates of teammate 1 in Control are also significantly inflated, at 41% for both genders. To bolster the claim of overconfidence, we need to compare beliefs across Main and Control: indeed for males and females beliefs about own performance are significantly higher than beliefs about another teammate 1 (Wilcoxon rank-sum test p-values 0.000 and 0.013 respectively). Overall we can say that overconfidence bias is strongly present for those who perform in the bottom half, and that it is significantly greater for males compared to females.

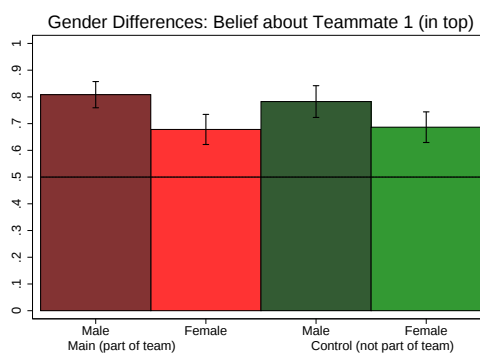
Finally Figure OA.6 presents prior beliefs, by gender, for the performance of teammate 2. Since, intuitively, there are no performance differences for teammate 2 by the subject's gender we do not resort to the sub-sample investigation in the previous figures. It suffices to report that there do not appear to be any gender differences in subject's beliefs about the performance of teammate 2, whether in Main or Control.

⁴Recall that in the Control treatment individuals observe the question responses of another different teammate 1. Thus they have additional information about performance of these individuals, explaining the large differences in reported beliefs, which are in the expected direction.

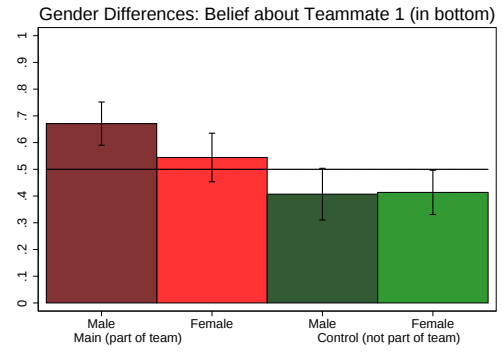
Figure OA.5: Gender Differences



(a) Full sample. Beliefs about teammate 1 being in the top half of performers.



(b) Sample restricted to teammate 1's that were in fact in the top half of performers (ranked 1-10 out of 20).

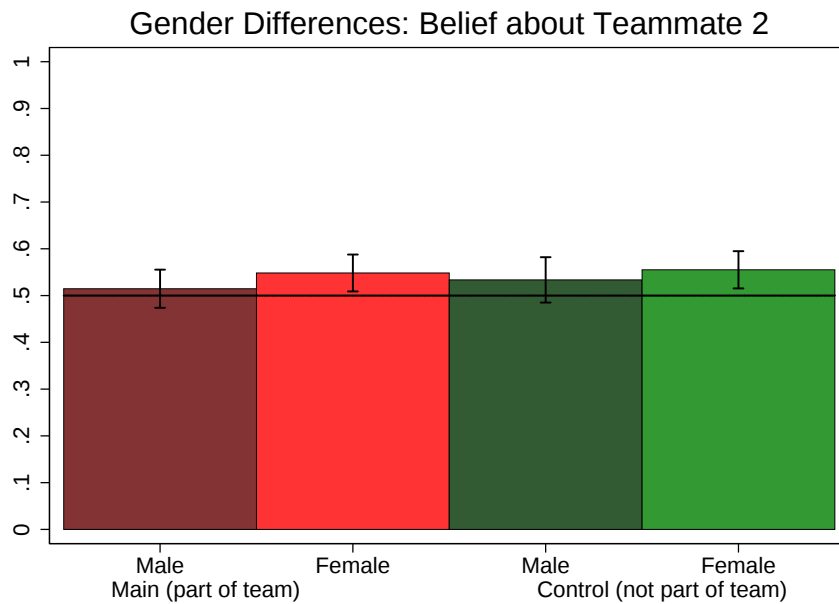


(c) Sample restricted to teammate 1's that were in fact in the bottom half of performers (ranked 11-20 out of 20).

3.2 Updating

After having noted significantly greater overconfidence for males relative to females, we now examine whether there are gender differences in updating these beliefs. Table OA.2 presents the main results (Table 1 in the main paper) for males and females separately. From these tables it is immediately apparent that the asymmetry in Main treatment appears to be driven by men. The positive asymmetry for men is large in magnitude, responsiveness to positive signals is 72% of a Bayesian, compared to negative signals which is 12% of a Bayesian, significant at the 1% level. In other words, men respond 6 times as much to positive feedback relative to negative feedback when it is about their own per-

Figure OA.6



Examining gender differences in initial beliefs about teammate 2, before receiving feedback. In Main treatment subject is on the team, as teammate 1. In Control, subject is not on the team. In both cases subject is aware of only the number of questions attempted by teammate 2.

formance. Examining the Control group, we find that in fact men respond slightly more to negative feedback when updating about others' performance, though the asymmetry is not significant. This observed asymmetry in Main is significantly different from that in Control at the 1% level (Chow-Test p-value 0.001). On the other hand, females exhibit very little asymmetry, and we are unable to reject that they respond equally to positive or negative signals. This is true both when they update about their own performance (Main) as well as when they update about another teammate 1's performance (Control). Consistent with some previous work, we see some evidence that females update more conservatively than males in the Control treatment. In Main they also react less strongly to positive signals, however, they show a similarly strong reaction to negative signals.

Table OA.1: Updating Beliefs about Teammate 1 (By Gender)

Regressor	Male		Female	
	Main (1)	Control (2)	Main (3)	Control (4)
δ	0.779** (0.086)	0.758*** (0.069)	0.682*** (0.073)	0.740*** (0.058)
β_1	0.722** (0.137)	0.552*** (0.117)	0.409*** (0.086)	0.512*** (0.102)
β_0	0.119*** (0.113)	0.671*** (0.112)	0.390*** (0.074)	0.390*** (0.069)
P-Value ($\delta = 1$)	0.0123	0.0008	0.0000	0.0000
P-Value ($\beta_1 = 1$)	0.0459	0.0003	0.0000	0.0000
P-Value ($\beta_0 = 1$)	0.0000	0.0043	0.0000	0.0000
P-Value ($\beta_1 = \beta_0$)	0.0072	0.4094	0.8796	0.3034
R^2	0.61	0.62	0.50	0.58
Observations	354	321	444	493
P-Value [Chow-test] for δ		0.8490		0.5290
P-Value [Chow-test] for β_1		0.3450		0.4371
P-Value [Chow-test] for β_0		0.0005		0.9994
P-Value [Chow-test] for $(\beta_1 - \beta_0)$		0.0053		0.5489

Analysis uses OLS regression. Difference is *significant from 1* at * 0.1; ** 0.05; *** 0.01. Robust standard errors clustered at individual level. R^2 corrected for no-constant.

Table OA.2 presents the analogue of Table 2 in the main paper, updating about teammate 2, for males and females separately. The only significant asymmetry arises for males evaluating the performance of their teammates when they themselves are a member of the team (Main). Males respond approximately 2 times as much to positive signals compared to negative signals, significant at the 5% level and echoing the direction of the overall results in the main paper. However, just as in the main paper, the asymmetry cannot be statistically distinguished from the asymmetry in the control group, in column 2. Regarding females, in columns 3 and 4 there does not appear asymmetry in response to feedback either in Main or Control.

Table OA.2: Updating Beliefs about Teammate 2 (By Gender)

Regressor	Male		Female	
	Main (1)	Control (2)	Main (3)	Control (4)
δ	0.843** (0.062)	0.608*** (0.082)	0.685*** (0.072)	0.832*** (0.046)
β_1	0.523*** (0.098)	0.605*** (0.126)	0.290*** (0.068)	0.411*** (0.074)
β_0	0.263*** (0.061)	0.459*** (0.116)	0.268*** (0.071)	0.393*** (0.066)
P-Value ($\delta = 1$)	0.0136	0.0000	0.0000	0.0004
P-Value ($\beta_1 = 1$)	0.0000	0.0023	0.0000	0.0000
P-Value ($\beta_0 = 1$)	0.0000	0.0000	0.0000	0.0000
P-Value ($\beta_1 = \beta_0$)	0.0384	0.2790	0.8115	0.8430
R^2	0.62	0.40	0.45	0.62
Observations	448	370	482	531
P-Value [Chow-test] for δ		0.0216		0.0837
P-Value [Chow-test] for β_1		0.6018		0.2258
P-Value [Chow-test] for β_0		0.1330		0.1949
P-Value [Chow-test] for $(\beta_1 - \beta_0)$		0.5333		0.9734

Analysis uses OLS regression. Difference is *significant from 1* at * 0.1; ** 0.05; *** 0.01. Robust standard errors clustered at individual level. R^2 corrected for no-constant.

4 Chosen Weights

While the primary focus of the empirical analysis is on determinants of beliefs and belief updating, it is informative to investigate how beliefs and updating affect subject's weighting decisions. Recall that individuals had to choose a weight from 0 to 1, with 0 representing all of the weight on teammate 2, and 1 representing all of the weight on teammate 1. Here we evaluate optimal weights relative to the Bayesian prediction: the weight chosen should be invariant to feedback. While feedback will impact beliefs, it does so proportionately for both teammates, leaving the weight unchanged. That is, after controlling for the initial weight, neither positive nor negative feedback should alter the submitted weight.

Table OA.3 shows regressions which examine impacts of subject characteristics and the main treatment on weighting decisions. The Bayesian prediction is that the initial weight in round one should have a coefficient of one, and all other coefficients should

be zero. From the table one can see that this is not the case. While the initial weight is positive and significant, it is less than one. What is more interesting is that against the Bayesian predictions, positive feedback has a statistically significant effect on the weight chosen, in columns (1) and (2). Additionally, there is some evidence that being a member of the team, i.e. our Main treatment, has a statistically significant effect on the chosen weight.

Yet, as columns (3) and (4) show, the positive effect of both a positive signal and the Main treatment are coming from the interaction between the two. In particular, this interaction increases the weight by 6.4 percentage points. This is about an 11% increase on the average weight chosen. Thus, when individuals are part of the team, when receiving a positive signal they increase the weight on their own performance by 6.4 percentage points, despite the Bayesian benchmark being to not alter the weight.

The result that there is some limited evidence of a larger weight after positive signals is consistent with the results on positively biased updating. Since subjects were also positively biased in updating about their teammate, this creates an overall moderating effect: the positive bias for both teammates works to cancel out, producing a more moderate weight report. A slight effect for positive signals is consistent with the slight overweighting of positive signals for self relative to teammate 2, while for negative signals there was no significant difference in the structural framework. In the Control treatment the responsiveness to feedback was balanced across both teammate 1 and 2, and for both positive and negative feedback. This is consistent with the results in Table [OA.3](#).⁵

⁵Finally, 7% of observations involved the submission of a different weight than what was recommended by z-tree. The average difference from the optimal recommended by z-tree was 0.056. However there are no systematic differences in submitting a different weight behavior by treatment.

Table OA.3: Submitted Weight on teammate 1

	(1)	(2)	(3)	(4)
Initial Weight	0.600*** (0.033)	0.515*** (0.042)	0.518*** (0.042)	0.473*** (0.044)
+ Signal	5.435*** (1.458)	5.364*** (1.435)	2.367 (2.136)	0.521 (2.081)
Main Treatment	3.540 (2.320)	4.267* (2.273)	1.210 (2.843)	0.854 (2.726)
+ Signal $\times$ Main Treatment			5.982** (2.833)	6.435** (2.738)
Female	2.767 (2.271)	2.223 (2.194)	2.480 (2.179)	2.244 (2.143)
Age	-0.387 (0.237)	-0.409* (0.241)	-0.410* (0.241)	-0.381 (0.236)
# Attempted by teammate 1		2.976*** (0.626)	2.965*** (0.626)	1.675** (0.687)
# Attempted by teammate 2		-1.272** (0.558)	-1.276** (0.555)	-1.675*** (0.528)
Score of teammate 1 on IQ Task				0.608*** (0.158)
Round Fixed Effects	✓	✓	✓	✓
R^2	0.38	0.40	0.40	0.42
Observations	2595	2595	2595	2595

Analysis uses OLS regression. Difference is significant from 0 at * 0.1; ** 0.05; *** 0.01. Robust standard errors clustered at individual level.

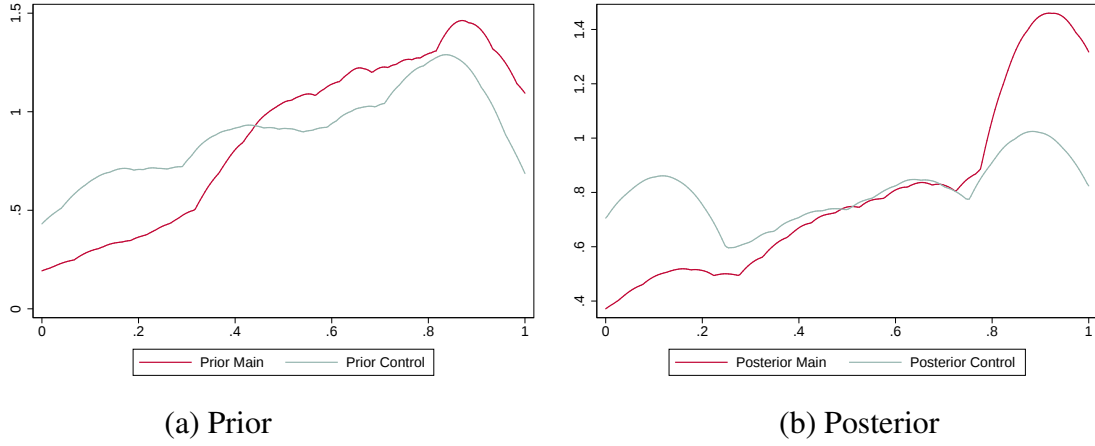
5 Examining Prior and Posterior Beliefs

Figures OA.7 and OA.8 present the distribution of beliefs for teammate 1 and teammate 2 respectively, separated into (a) prior and (b) posterior, by treatment. We reject that beliefs are normally distributed for all sub-samples at the 1% level, with the exception of prior beliefs in Control about teammate 2 which we reject at the 10% level (Shapiro-Wilk test).

Regarding teammate 1, relative to the Control treatment, subjects in the Main treatment are initially more overconfident (significantly different at 1% level; Kolmogorov-Smirnov test). After four rounds of feedback, the belief distributions diverge even further, with Main subjects becoming more overconfident, and the distributions continue to be significantly different at the 1% level. Regarding teammate 2, there are no prior differences in confidence, and the distributions are not significantly different at conventional

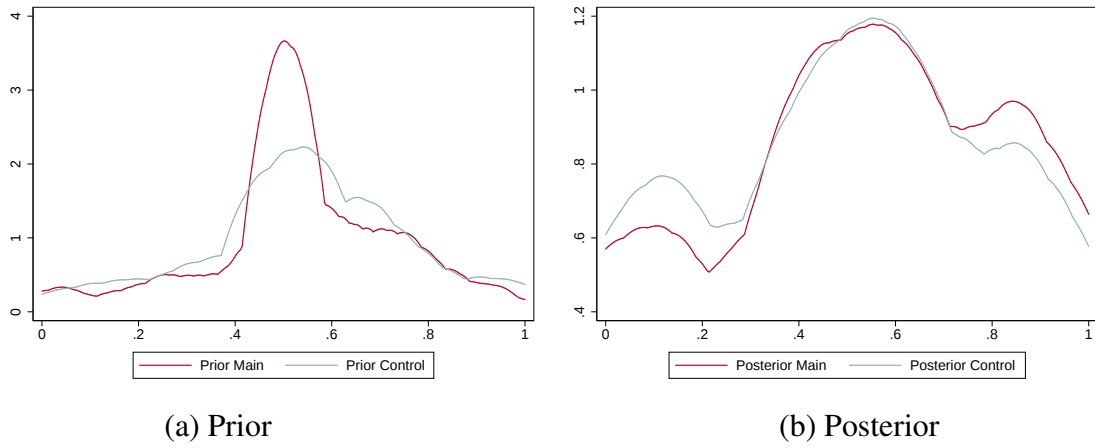
levels (Kolmogorov-Smirnov test). After four rounds of feedback, there is a shift towards greater confidence in Main relative to Control, although the distributions continue to not be significantly different.

Figure OA.7: Prior and Posterior Belief Distribution for Teammate 1



Epanechnikov kernel density estimate. (b) Posterior is calculated after receiving four signals.

Figure OA.8: Prior and Posterior Belief Distribution for Teammate 2



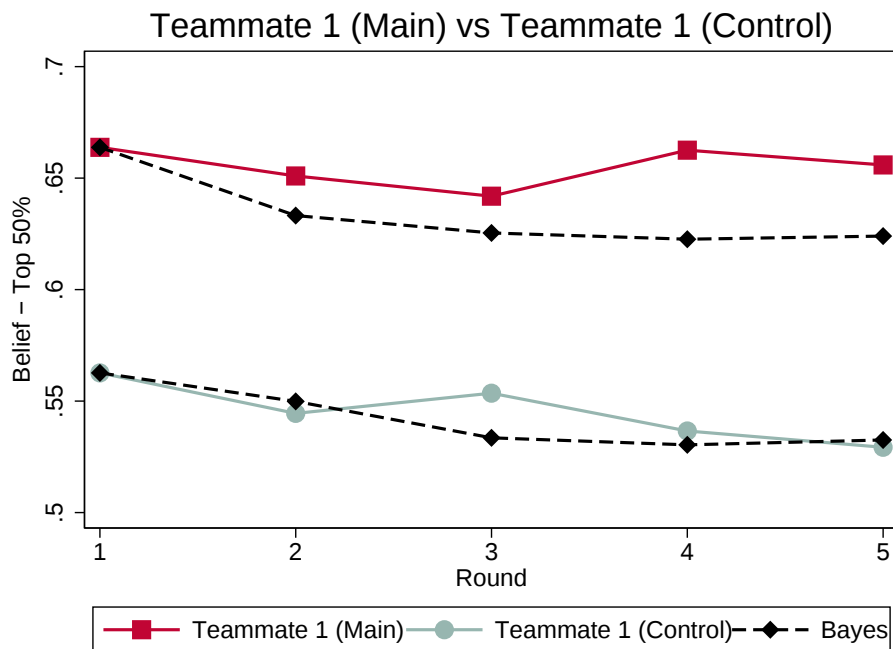
Epanechnikov kernel density estimate. (b) Posterior is calculated after receiving four signals.

Figures OA.9 and OA.10 examine the evolution of beliefs in response to feedback for teammate 1 and 2 respectively, starting from the first prior, before receiving any feedback. While posterior beliefs about one's self (Main, teammate 1) are significantly greater than beliefs about teammate 1 in the Control, this is in large part driven by differences in prior beliefs due to overconfidence. In both figures one can see a pattern that posterior beliefs in

the final round deviate further from the Bayesian prediction in Main compared to Control, both for teammate 1 and 2.

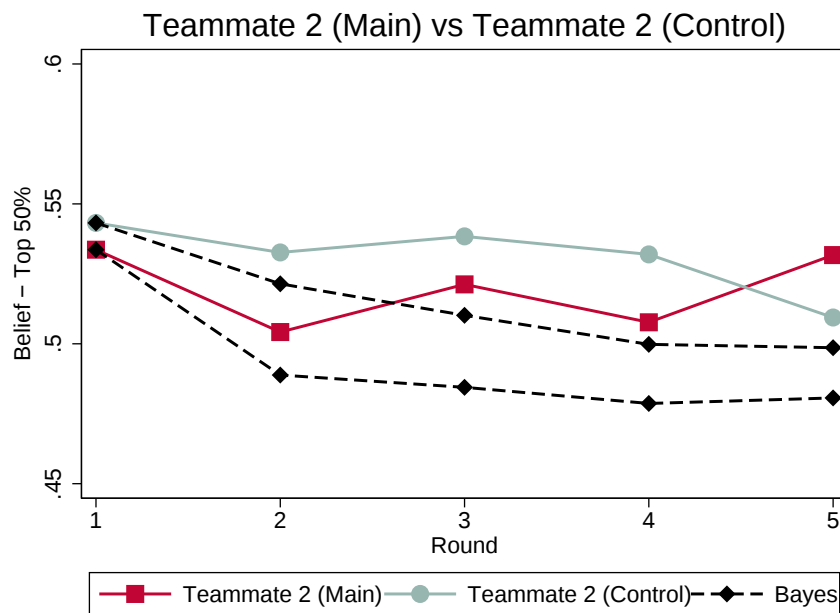
Figure OA.11 examines this more closely, presenting the difference between reported posteriors and the Bayesian prediction given subjects' initial priors, after four rounds of feedback. This corresponds to round 5 in the two figures above. While this does present evidence that positive deviations are more pronounced in the Main treatment, we also note that the difference between the deviations in Main and Control are not significantly different at conventional levels.

Figure OA.9: Evolution of Beliefs: teammate 1



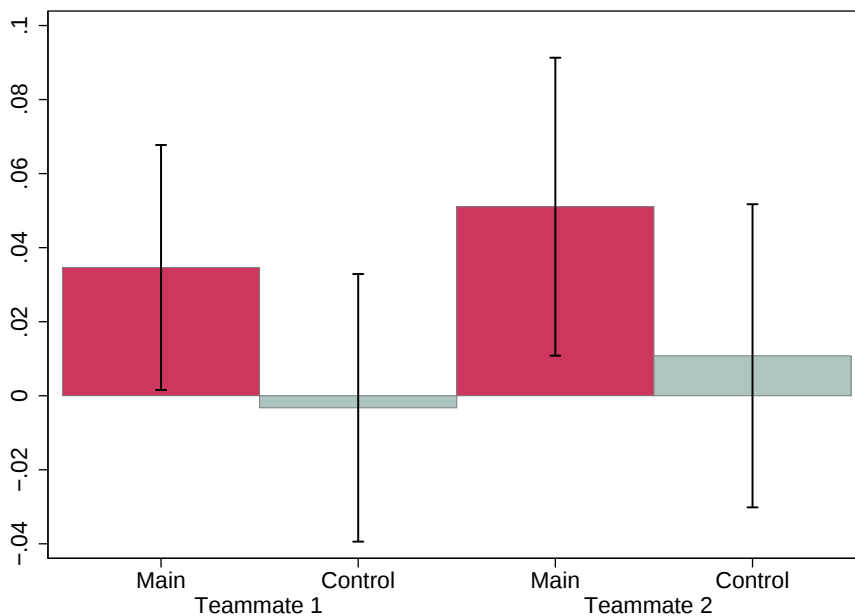
Evolution of beliefs about teammate 1 starting from prior beliefs with 4 round of feedback. Bayesian benchmark is calculated from subject's first prior, then evolves given actual signals observed. Standard error bars omitted for clarity (error bars are always overlapping with bayesian predictions).

Figure OA.10: Evolution of Beliefs: teammate 2



Evolution of beliefs about teammate 2 starting from prior beliefs with 4 round of feedback. Bayesian benchmark is calculated from subject's first prior, then evolves given actual signals observed. Standard error bars omitted for clarity (error bars are always overlapping with bayesian predictions).

Figure OA.11: Raw Deviation of Posterior Beliefs from Bayesian Benchmark



Plot of the difference between Posterior beliefs and Bayesian beliefs after four rounds of feedback. Bayesian beliefs are calculated using subject priors before any feedback.

6 Sampling Robustness Checks

6.1 Excluding Part 3 from analysis

At the end of Part 2, subjects had the opportunity to re-match with new teammate 2's. In Part 3, those who are not re-matched, continue to update about the same teammate 2, while those who are re-matched state a prior belief about the new teammate 2, and then update about this new teammate. In the statistical analysis of the main paper these situations are pooled together. Here we replicate the main analysis, the analogous tables for Table 1 and 2 in the main paper, excluding Part 3 of the experiment. Table OA.4 presents this analysis for teammate 1, while Table OA.5 presents the analysis for teammate 2. From these tables it is apparent that the results are not significantly affected by the removal of Part 3.

Table OA.4: Updating Beliefs about Teammate 1

Regressor	(1) Main Treatment	(2) Control Treatment
δ	0.702*** (0.060)	0.709*** (0.053)
β_1	0.671*** (0.074)	0.581*** (0.095)
β_0	0.325*** (0.066)	0.609*** (0.082)
P-Value ($\delta = 1$)	0.0000	0.0000
P-Value ($\beta_1 = 1$)	0.0000	0.0000
P-Value ($\beta_0 = 1$)	0.0000	0.0000
P-Value ($\beta_1 = \beta_0$)	0.0018	0.8032
R^2	0.53	0.55
Observations	569	525
P-Value [Chow-test] for δ (Regressions (1) and (2))	0.9325	
P-Value [Chow-test] for β_1 (Regressions (1) and (2))	0.4524	
P-Value [Chow-test] for β_0 (Regressions (1) and (2))	0.0067	
P-Value [Chow-test] for $(\beta_1 - \beta_0)$ (Regressions (1) and (2))	0.0163	

Analysis uses OLS regression, for only Part 2 of the experiment. Difference is *significant from 1* at * 0.1; ** 0.05; *** 0.01. Robust standard errors clustered at individual level. R^2 corrected for no-constant. δ is the coefficient on the log prior odds ratio. β_1 and β_0 are coefficients on the log likelihood of observing positive and negative signals respectively. Constant omitted because of collinearity. Bayesian updating corresponds to $\delta = \beta_1 = \beta_0 = 1$. $\beta_1, \beta_0 < 1$ indicates conservative updating. $\beta_1 - \beta_0 > 0$ indicates positive asymmetric updating.

Table OA.5: Updating Beliefs about Teammate 2

Regressor	(1) Main Treatment	(2) Control Treatment
δ	0.690*** (0.062)	0.664*** (0.062)
β_1	0.405*** (0.064)	0.523*** (0.080)
β_0	0.238*** (0.058)	0.419*** (0.076)
P-Value ($\delta = 1$)	0.0000	0.0000
P-Value ($\beta_1 = 1$)	0.0000	0.0000
P-Value ($\beta_0 = 1$)	0.0000	0.0000
P-Value ($\beta_1 = \beta_0$)	0.0597	0.2914
R^2	0.39	0.40
Observations	672	582
P-Value [Chow-test] for δ (Regressions (1) and (2))		0.7704
P-Value [Chow-test] for β_1 (Regressions (1) and (2))		0.2508
P-Value [Chow-test] for β_0 (Regressions (1) and (2))		0.0566
P-Value [Chow-test] for $(\beta_1 - \beta_0)$ (Regressions (1) and (2))		0.6250

Analysis uses OLS regression, for only Part 2 of the experiment. Difference is *significant from 1* at * 0.1; ** 0.05; *** 0.01. Robust standard errors clustered at individual level. R^2 corrected for no-constant. δ is the coefficient on the log prior odds ratio. β_1 and β_0 are coefficients on the log likelihood of observing positive and negative signals respectively. Constant omitted because of collinearity. Bayesian updating corresponds to $\delta = \beta_1 = \beta_0 = 1$. $\beta_1, \beta_0 < 1$ indicates conservative updating. $\beta_1 - \beta_0 > 0$ indicates positive asymmetric updating.

6.2 Sub-sample robustness checks

Here we will change some of the inclusion criteria used in Tables 1 and 2 in the main paper. In these tables we had excluded wrong updates as well as boundary observations (beliefs of 0 or 1). An update is defined as being in the wrong direction when individuals update at least one of the two beliefs per feedback round, for either teammate 1 or 2, in the opposite direction than feedback suggests, *without* adjusting the other belief in the direction feedback suggests. Another restriction we use in the main paper is to exclude boundary observations of beliefs equal to either 0 or 1, since these drop out of the logit function.

Table OA.6 presents the implications of these sampling restrictions for analyzing reactions to feedback of teammate 1. In Column 1 of Table OA.6 we include observations which involve updating in the wrong direction following our earlier definition. One can

see that while the magnitudes are affected, the patterns are virtually unchanged. If anything there is slightly more evidence that positive signals receive more weight in Main (beliefs about own performance) relative to Control (beliefs about another's performance). In Column 2 we include boundary observations, noting that this as well does not significantly alter the main pattern of results, though it does create some asymmetry in Control.⁶ Finally in Column 3 we remove all restrictions to the data. The patterns are similar, though again with some asymmetry in Control. We note that the overall results appear to be even stronger than our main analysis suggests, pointing to a stronger reaction to positive signals (β_1) in Main, a weaker reaction to negative signals (β_0), and a rejection that the size of the asymmetry is equal between Main and Control.

Similarly, in Table [OA.7](#) we investigate the implications of removing these sampling restrictions on updating about teammate 2. We see similar patterns, with some stronger evidence that positive signals are incorporated less in Main (subject is part of the team) rather than control, compared to Table 2 in the paper. Overall these tables suggest that our patterns of results are consistent for other sampling strategies.

⁶We do not have an explanation for this asymmetry, but we note that including boundary observations essentially amounts to adding noise to the model, since once beliefs arrive at the extremes, they typically stay there. Regardless, we are still nearly able to reject that the asymmetry in Main is greater than Control, as the p-value is 0.115.

Table OA.6: Updating Beliefs about Teammate 1

Regressor	Inc. Wrong Dir.		Inc. Boundary		Inc. All	
	Main (1)	Control (2)	Main (3)	Control (4)	Main (5)	Control (6)
δ	0.703*** (0.056)	0.710*** (0.048)	0.826*** (0.048)	0.842*** (0.044)	0.790*** (0.052)	0.805*** (0.046)
β_1	0.513*** (0.072)	0.365*** (0.070)	0.743*** (0.084)	0.646*** (0.084)	0.658*** (0.083)	0.467*** (0.080)
β_0	0.071*** (0.058)	0.339*** (0.053)	0.281*** (0.075)	0.445*** (0.078)	0.054*** (0.069)	0.258*** (0.067)
P-Value ($\delta = 1$)	0.0000	0.0000	0.0004	0.0005	0.0001	0.0000
P-Value ($\beta_1 = 1$)	0.0000	0.0000	0.0024	0.0000	0.0001	0.0000
P-Value ($\beta_0 = 1$)	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000
P-Value ($\beta_1 = \beta_0$)	0.0001	0.7577	0.0003	0.0653	0.0000	0.0416
R^2	0.50	0.54	0.55	0.57	0.49	0.51
Observations	996	985	938	899	1077	1062
P-Value [Chow-test] for:						
δ		0.9210		0.8037		0.8263
β_1		0.1408		0.4127		0.0959
β_0		0.0006		0.1297		0.0342
($\beta_1 - \beta_0$)		0.0026		0.1153		0.0141

Analysis uses OLS regression. Difference is *significant from 1* at * 0.1; ** 0.05; *** 0.01. Robust standard errors clustered at individual level. R^2 corrected for no-constant. δ is the coefficient on the log prior odds ratio. β_1 and β_0 are coefficients on the log likelihood of observing positive and negative signals respectively. Constant omitted because of collinearity. Bayesian updating corresponds to $\delta = \beta_1 = \beta_0 = 1$. $\beta_1, \beta_0 < 1$ indicates conservative updating. $\beta_1 - \beta_0 > 0$ indicates positive asymmetric updating.

Table OA.7: Updating Beliefs about Teammate 2

Regressor	Inc. Wrong Dir.		Inc. Boundary		Inc. All	
	Main (1)	Control (2)	Main (3)	Control (4)	Main (5)	Control (6)
δ	0.741*** (0.044)	0.687*** (0.047)	0.868*** (0.044)	0.761*** (0.054)	0.845*** (0.040)	0.733*** (0.051)
β_1	0.330*** (0.054)	0.342*** (0.066)	0.435*** (0.073)	0.666*** (0.080)	0.330*** (0.069)	0.497*** (0.075)
β_0	0.133*** (0.038)	0.275*** (0.056)	0.389*** (0.066)	0.549*** (0.071)	0.239*** (0.059)	0.379*** (0.065)
P-Value ($\delta = 1$)	0.0000	0.0000	0.0030	0.0000	0.0002	0.0000
P-Value ($\beta_1 = 1$)	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000
P-Value ($\beta_0 = 1$)	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000
P-Value ($\beta_1 = \beta_0$)	0.0042	0.4056	0.6331	0.1979	0.3120	0.1903
R^2	0.51	0.45	0.47	0.45	0.45	0.41
Observations	1166	1088	1085	974	1244	1150
P-Value [Chow-test] for:						
δ		0.4052		0.1231		0.0838
β_1		0.8942		0.0316		0.0995
β_0		0.0367		0.0993		0.1103
($\beta_1 - \beta_0$)		0.2163		0.5880		0.8273

Analysis uses OLS regression. Difference is *significant from 1* at * 0.1; ** 0.05; *** 0.01. Robust standard errors clustered at individual level. R^2 corrected for no-constant. δ is the coefficient on the log prior odds ratio. β_1 and β_0 are coefficients on the log likelihood of observing positive and negative signals respectively. Constant omitted because of collinearity. Bayesian updating corresponds to $\delta = \beta_1 = \beta_0 = 1$. $\beta_1, \beta_0 < 1$ indicates conservative updating. $\beta_1 - \beta_0 > 0$ indicates positive asymmetric updating.

7 Alternative Explanations

7.1 Anchoring

One potential alternative explanation for our results is that there is some form of anchoring bias which causes individuals to update in similar ways about their teammate, in a purely mechanical (and naive) sense. To examine the scope for this, we examine the raw correlations among beliefs about teammate 1 and 2. First we find that the correlation between raw beliefs for teammate 1 and 2 is low, with a correlation of $\rho = 0.058$ in the Main treatment (compared with -0.003) in the Control treatment.⁷ However, it may be even more important to examine the change in beliefs. In order to account for possible heuristical updating which may lead to similar updating patterns across teammate 1 and 2, we examine the correlation between the absolute change in beliefs ($posterior - prior$) for teammate 1 versus teammate 2. We also examine the correlation between the percentage change in beliefs ($\frac{posterior - prior}{prior}$). Table OA.8 shows these correlation for Main versus Control, split by signal type, as well as overall, pooled over all rounds.

Table OA.8: Correlation in Belief Updates across Teammate 1 and 2

Signal	Main			Control		
	+	−	All	+	−	All
Absolute	−0.059	−0.185	−0.125	0.021	−0.085	−0.018
Percentage	0.032	0.004	0.008	0.068	0.219	0.159
<i>N</i>	695	691	1386	641	648	1289

Table presents the correlation across teammate 1 and 2 of the absolute difference, $posterior - prior$ in (Row 1), or the percentage difference, $\frac{posterior - prior}{prior}$ in (Row 2).

From Table OA.8 one can see in fact that raw movements in belief updates are not significantly positively correlated across teammate 1 and 2 in the Main treatment. In fact, a negative correlation emerges when considering absolute difference updating. Rather, it is in the Control group where a positive correlation (for percentage updating) can be found. Thus we conclude that it is unlikely that simple anchoring heuristics, such as increasing beliefs by the same fixed amount or same fixed percentage for both teammates are driving the patterns of results.

⁷This correlation in the Main treatment is much lower for Part 2 only, at 0.012, where significant asymmetric updating is also found, see Online Appendix Table 1.

7.2 In-Group Bias

We also briefly consider the possibility that individuals may become overconfident or optimistic about team performance, rather than individual performance. First we note that we don't find any bias when individuals update about a team in our Control treatment, so such a theory would nonetheless require that the individual herself be a part of the team.

Next we don't find any initial overestimation of teammate 2's performance in Main compared to teammate 2's performance in Control – the belief that teammate 2 scores in the top 50% are 53.4% and 54.3% (Wilcoxon rank-sum test p-value: 0.572) in the Main and Control treatment, respectively. Thus, the in-group bias would need to selectively apply to belief updating from signals, but not appear in initial belief formation. For these reasons, while we cannot completely rule out a role for such a team bias in our results, the above patterns suggest it is unlikely to solely be able to account for our findings.

7.3 Selective Discounting or Ignorance

Here we consider a possible alternative explanation which involves selectively discounting or ignoring particular signals. First we briefly consider a hypothetical individual who discounts the true signal strength of negative feedback overall, independently of whether they are updating about themselves (teammate 1) or their teammate (teammate 2). We note that such a model would predict *equivalent* asymmetry across self and teammate. However, examining the first columns in Tables 2 and 3 in the main paper, we note that we rejected this prediction, at the 10% level.

Next we turn to selective ignorance of negative signals. To set up this explanation, we consider a different variation of our theoretical model where the sub-conscious process can selectively choose whether to ignore new information. That is, the sub-conscious process chooses the optimal action of whether or not to update, conditional on observing the signal. In this model the action is binary, update or not. We assume otherwise updating would follow Bayes' rule.

We note that it would always be optimal to update after a positive signal, since the benefits from overconfidence and the benefits from accuracy go in the same direction. It may or may not be optimal to update after a negative signal, since the costs to holding lower beliefs are in opposition to the benefits of accuracy. Finally, we note that adding a cognitive fixed cost of updating, would potentially lower the number of updates, but not in a biased way between positive versus negative signals.

To test the predictions of such a model we investigate whether the propensity to update differed across our Main and Control treatments. In fact in the Control treatment the

probability of updating after any signal is 67%, while in Main it is 63%. The difference is not significant when accounting for potential correlation of standard errors at the individual level. Moreover, this slight difference is driven entirely by positive signals, not negative signals. Thus there is no evidence that subjects ignore negative feedback more in the Main treatment. Beyond this, in fact subjects are more likely to update after a negative signal than a positive signal (in both Main and Control), which further goes against this potential theoretical explanation.

8 Alternative Models

Here we present two alternative models of self-serving attribution bias which focus on two relatively stark predictions involving self-serving beliefs which are formed due to mis-attribution towards either (i) noise, or (ii) the external fundamental (teammate 2). The first we denote by: noisy attribution bias (NAB). With NAB the individual processes information about other factors (teammate 2) accurately, but is positively biased about own performance (teammate 1) at the expense of noise. In the second model, fundamental attribution bias (FAB), the individual respects the amount of noise contained in the signal, but is biased about her own performance (teammate 1) at the expense of teammate 2.

8.1 Updating with Noisy Attribution Bias

With NAB, individuals over-attribute positive feedback to their own performance, and under-attribute negative feedback to bad luck. We additionally specify that NAB predicts that individuals update using Bayes' rule regarding the performance of this teammate. Someone who exhibits NAB will update in a way that is consistent with mis-interpreting the strength of the binary signal. That is, when they receive a positive signal, they believe it is more informative about their performance than it really is. When they receive a negative signal, they believe it is less informative about their performance than it really is. They over-interpret the strength of the signal by a factor of γ_p , where $\gamma_p \geq 1$ in the case of a positive signal, and $\gamma_n \geq 1$ in the case of a negative signal. However, updating about teammate 2's performance occurs as if the strength of the signal were correctly interpreted in all states, i.e. updating about teammate 2 is Bayesian.

Thus, regarding own performance, and following the notation in the main paper, biased updating in response to positive and negative feedback through NAB results in up-

ward biased beliefs:

$$[b_{t+1}^{1,NAB}|s_t = p] = \frac{\gamma_p [\Phi_{TT}b_t^{TT} + \Phi_{TB}b_t^{TB}]}{\gamma_p [\Phi_{TT}b_t^{TT} + \Phi_{TB}b_t^{TB}] + \Phi_{BT}b_t^{BT} + \Phi_{BB}b_t^{BB}} \geq [b_{t+1}^{1,BAYES}|s_t = p], \quad (1)$$

$$[b_{t+1}^{1,NAB}|s_t = n] = \frac{\gamma_n [(1 - \Phi_{TT})b_t^{TT} + (1 - \Phi_{TB})b_t^{TB}]}{\gamma_n [(1 - \Phi_{TT})b_t^{TT} + (1 - \Phi_{TB})b_t^{TB}] + (1 - \Phi_{BT})b_t^{BT} + (1 - \Phi_{BB})b_t^{BB}} \quad (2)$$

$$\geq [b_{t+1}^{1,BAYES}|s_t = n].$$

With NAB, updating about their own performance exhibits positive asymmetry - individuals over-weight positive signals and under-weight negative signals relative to a Bayesian. They are perfectly Bayesian with regards to their teammate's performance.

8.2 Updating with Fundamental Attribution Bias

With FAB, individuals over-attribute positive feedback to their own performance, at the expense of the other source of uncertainty, i.e. their teammate. Similarly, they under-attribute negative feedback to themselves, and over-attribute it to their teammate.

FAB takes the same functional form as NAB with regards to own performance.

$$[b_{t+1}^{1,FAB}|s_t = p] = [b_{t+1}^{1,NAB}|s_t = p] = \quad (3)$$

$$\frac{\gamma_p [\Phi_{TT}b_t^{TT} + \Phi_{TB}b_t^{TB}]}{\gamma_p [\Phi_{TT}b_t^{TT} + \Phi_{TB}b_t^{TB}] + \Phi_{BT}b_t^{BT} + \Phi_{BB}b_t^{BB}} \geq [b_{t+1}^{1,BAYES}|s_t = p],$$

$$[b_{t+1}^{1,FAB}|s_t = n] = [b_{t+1}^{1,NAB}|s_t = n] = \quad (4)$$

$$\frac{\gamma_n [(1 - \Phi_{TT})b_t^{TT} + (1 - \Phi_{TB})b_t^{TB}]}{\gamma_n [(1 - \Phi_{TT})b_t^{TT} + (1 - \Phi_{TB})b_t^{TB}] + (1 - \Phi_{BT})b_t^{BT} + (1 - \Phi_{BB})b_t^{BB}} \geq [b_{t+1}^{1,BAYES}|s_t = n].$$

With FAB the individual updates in a biased but consistent manner across themselves and their teammate, unlike with NAB, where they update as a standard Bayesian with respect to their teammate. In response to positive and negative signals respectively FAB implies for teammate 2:

$$[\hat{b}_{t+1}^2 | s_t = p] = \frac{\gamma_p \Phi_{TT} b_t^{TT} + \Phi_{BT} b_t^{BT}}{\gamma_p [\Phi_{TT} b_t^{TT} + \Phi_{TB} b_t^{TB}] + \Phi_{BT} b_t^{BT} + \Phi_{BB} b_t^{BB}} \quad (5)$$

$$[\hat{b}_{t+1}^2 | s_t = n] = \frac{\gamma_n (1 - \Phi_{TT}) b_t^{TT} + (1 - \Phi_{BT}) b_t^{BT}}{\gamma_n [(1 - \Phi_{TT}) b_t^{TT} + (1 - \Phi_{TB}) b_t^{TB}] + (1 - \Phi_{BT}) b_t^{BT} + (1 - \Phi_{BB}) b_t^{BB}} \quad (6)$$

Because $\Theta = \Phi_{TT}\Phi_{BB} - \Phi_{TB}\Phi_{BT} \leq 0$, this guarantees that $[b_{t+1}^{2,FAB} | s_t = s] \leq [b_{t+1}^{2,BAYES} | s_t = s]$. Thus, FAB implies that when considering their teammate's performance they update asymmetrically in the negative direction, i.e. they under (over)-weight positive (negative) signals.

We note that both NAB and FAB predict positive asymmetry with respect to own performance beliefs, but either (i) no asymmetry or (ii) negative asymmetry, respectively, regarding updating about the performance of teammate 2. As our main results in the paper show, neither of these two theoretical predictions are borne out in the data.

9 Instructions

As we have four sets of instructions (two waves $\times$ two versions (Main and Control)) we present here two for brevity: wave 1 for the Main treatment, and wave 2 for the Control treatment. Note that the differences between wave 1 and 2 are that wave 2 included an additional Part 3 where subjects had their willingness to pay (WTP) elicited to be matched to a new teammate 2. Part 3 then repeated Part 2 (updating with feedback) with either the old or new teammate, depending on the outcome of the WTP procedure. An additional difference between waves 1 and 2 is that wave 2 elicited beliefs about the full distribution (four states), while wave 1 elicited only beliefs about each teammate being in the top half.

The difference between Main and Control was that in Main the subjects themselves were playing the role of teammate 1, and hence their test performance was directly relevant for payoffs. In Control the subjects were making identical decisions for randomly selected teammates 1 and 2. Hence in Control the subjects own performance was not relevant.

9.1 Wave 1: Main Treatment

Welcome to this experiment!

In this experiment you will be asked to make several decisions that determine your earnings. For showing up on time you receive 5 EUR. During the experiment you can earn additional money. How much? That depends on the decisions that you make during the experiment. Please read and follow the instructions carefully. They contain everything you need to know to participate. All decisions are taken anonymously, that is, the identity of the participants is not revealed to the other participants at any point. Also the payout is going to take place anonymously at the end of the experiment.

Please note that from now on and during the entire experiment, no communication is permitted. Throughout the experiment the use of mobile phones, smartphones, tablets or the like is prohibited. An infringement of these rules will result in exclusion from the experiment and all payments.

If you have questions at any point, please stretch your hand out of the booth and the experimenter will come to you to answer them in private.

The experiment consists of 2 parts. In the following we will explain to you Part 1 of the experiment. We will present you with the instructions for Part 2 after you have finished Part 1.

Part 1 – The knowledge and IQ test

The computer screen will present you with a set of **15 trivia and logic questions**. Questions similar to these are often used to measure a person's general intelligence (IQ). Your task is to answer as many of these questions correctly as possible.

Your score is determined as follows. For every correct answer you receive 2.5 points, for every wrong answer you lose 1 point (you cannot end up with less than 0 points). Questions that you do not answer will not affect your score. Your score will be transformed into Euros at the end of the experiment. The exchange rate is 1 point = 0.10 EUR.

Payoff example:

Suppose you attempted 12 out of the 15 trivia and logic questions and answered 9 of them correctly and the remaining 3 incorrectly.

That means, your score is $9 \cdot 2.5 - 3 \cdot 1 = 19.5$ points

Your payments from this part of the experiment then amount to 1.95 EUR.

You will be given **10 minutes** to solve all questions. A timer in the upper right corner of your screen will inform you how much time (in seconds) is left to finish the test. You may use the backside of these instructions to take notes, if you wish to do so.

You will not learn your score until the end of the experiment.

Attention! The test will begin shortly on the screen!

Part 2 – The team task

There are no more tests for the remainder of this experiment.

In this part of the experiment you will be matched randomly with a teammate. You will have the chance to earn additional 10 EUR. How likely this is will depend on your score and your teammate's score in the previous test and the decisions that you are going to take in this part of the experiment.

Remember, the score depends on the number of correctly answered questions (+2.5 points each) and the number of incorrectly answered questions (-1 point each).

The identity of the teammate will never be revealed to you. The only information that the computer is going to present you with is the number of questions that you and the teammate *attempted* to answer in the test.

Before the start of Part 2 the computer has separately compared your and your teammate's scores with the *same* randomly selected group of 19 other participants in today's experiment:

- **You** have been compared with the 19 participants and are either **ranked in the Top 10 or Bottom 10** of the 20 performances.
- Similarly, **your teammate** has been compared with the same 19 participants and is similarly **ranked in either the Top 10 or Bottom 10**.

What determines whether you earn the additional 10 EUR in this part of the experiment?

- If both your and your teammate's scores on the test were in the **Top 10**, you will earn the 10 EUR for sure (100% probability of winning).
- If both your and your teammate's scores on the test were in the **Bottom 10**, you will NOT earn the 10 EUR for sure (0% probability of winning).
- If **ONE** of either your score or your teammate's score was in the Top 10, and the other was in the Bottom 10, **your probability of earning the 10 EUR depends on you choosing a "weight" from 0 to 100** that you assign to your own performance in the team.

Being in the Top 10 indicates that one is in the top 50% (half) of performers. Being in the Bottom 10 indicates that one is in the bottom 50% (half) of performers.

We will now explain this "weight" in more detail.

Part 2 – Choosing a weight

If ONE of either your score or your teammate's score was in the Top 10, and the other was in the Bottom 10, the probability of earning the 10 EUR is calculated as follows:

- a) If you scored in the Top 10 (and your teammate did not): $\sqrt{\text{Weight}/100}$
- b) If your teammate scored in the Top 10 (and you did not): $\sqrt{(100 - \text{Weight})/100}$.

Don't worry if this looks confusing. **You will be asked to enter the probabilities you believe you and your teammate scored in the Top 10.** Then the computer will automatically calculate the weight on your own performance that gives you the highest probability of earning the 10 EUR.

To guarantee the highest probability of earning the 10 EUR, you should choose a weight **closer to 100 if you believe you were more likely than your teammate to score in the Top 10**, and **closer to 0 if you believe your teammate**

was more likely to score in the Top 10. If you think that both you and your teammate have the same likelihood of scoring in the Top 10, then the optimal weight would be 50.

Here’s an example of how it works. Suppose you believe that there is a 25% chance your performance was in the Top 10, but you think there’s a 50% chance your teammate’s performance was in the Top 10. The computer will calculate the probability of earning the 10 EUR for every possible weight from 0 to 100 and inform you which weight you should *optimally* choose given the probabilities that you enter.

Given the above probabilities, this is how the information would be presented to you on the screen:

Enter your estimates for the probability your performance was in the Top 10, and for the probability your teammates performance was in the Top 10.

In order to have the greatest chance of earning the 10 EUR, you should enter these probabilities as accurately as possible.

Probability your performance was in the Top 10:

25 %

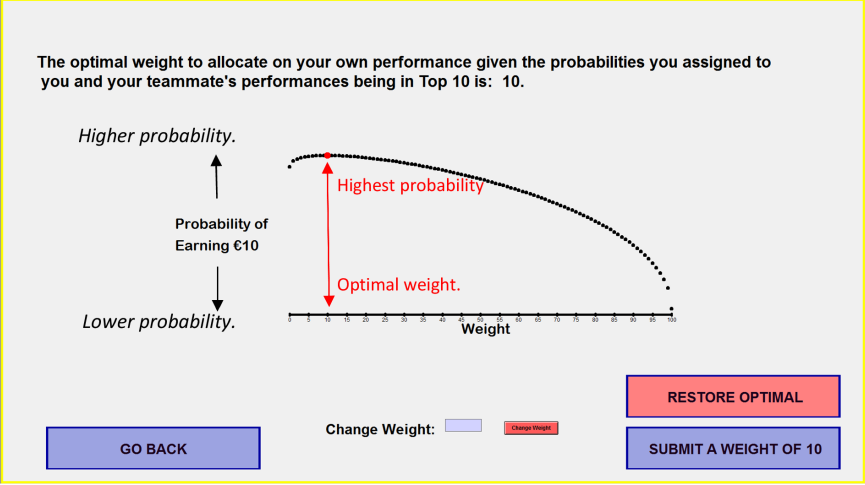
Probability your teammate's performance was in the Top 10:

50 %

Click the slider to adjust the probabilities.

CALCULATE OPTIMAL WEIGHT

(You will be able to go back and change your estimates)



Summary

In order to have the best chance of earning the 10 EUR, it pays to estimate your performance and your teammate's performance as accurately as possible. The computer then calculates the optimal weight for you.

If you are satisfied with this weight, you can "confirm" this weight on the following screen and submit it for evaluation.

OR

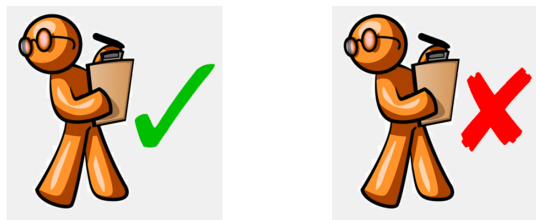
If you want to put a different weight, you may either do this directly,

OR

You may go back to the previous screen and enter different probabilities with which you expect that your and the teammate's performances were in the Top 10. Based on this, a new weight will be calculated and displayed to you.

Feedback

After you will have submitted the weight, on the next screen, a Team Evaluator will give you either a "Green Check" or a "Red X", which is related to your performance and your teammate's performance.



However note that Team Evaluators can sometimes make mistakes:

- If you both scored in the Top 10, a Team Evaluator will give you a Green Check in 9 out of 10 cases. In 1 out of 10 cases he will give you a Red X.
- If one of you scored Top 10 and the other one scored in the Bottom 10, a Team Evaluator will give you a Green Check in 5 out of 10 cases. In other 5 out of 10 cases, he will give you a Red X.
- If both of you scored in the Bottom 10, a Team Evaluator will give you a Green Check in 1 out of 10 cases. In 9 out of 10 cases he will give you a Red X.

A **Green Check** indicates that it is **unlikely** that **both of you** scored in the **Bottom 10**, while a **Red X** indicates that it is **unlikely** that **both of you** scored in the **Top 10**.

That means, while a team evaluator gives you some information about you and your teammate's performances, you do not learn your or your teammate's performances for sure.

What happens next?

After receiving the Team Evaluator's feedback the computer will ask you once more to **enter the probabilities you believe you and your teammate scored in the Top 10**. Then the computer will again automatically calculate the weight that gives you the highest probability of earning the 10 EUR.

As before, you may plug in and try out different numbers. Once you are satisfied with the weight that the computer calculates for you – or, alternatively, once you have entered a weight yourself – you can “confirm” this weight and submit it for evaluation.

This procedure (evaluator feedback – entering performance expectations and assigning weights) will be repeated 4 times. So that, in total, you receive feedback from 4 different evaluators, after which the computer asks you to enter your beliefs and weights again. In each round you have the chance to earn the 10 EUR in the way described above.

Please note: Each time you receive feedback from a Team Evaluator, it is *completely independent* of the feedback you previously received. Also note that the weight(s) you have chosen have no effect on feedback. Feedback is just about you and your teammate's performances.

What determines your earnings today?

Besides your earnings in the trivia and logic test in Part 1 of the experiment, ***one out of the five*** weighting decisions that you enter in Part 2 (your initial one and the four that you enter after having received evaluator feedback) will be randomly selected for payment. **Thus, every decision you make is important as it may affect your earnings if it is chosen for payment.**

At the end of the experiment the computer will inform you about its selection on your screen.

On the next page, we will ask you to answer some control questions. This is to make sure that all participants understand the instructions of this experiment. Once all participants have answered all questions correctly, Part 2 of the experiment will start.

If you have questions about any of the instructions (now or later), please stretch your hand out of the booth and the experimenter will come to you to answer them in private. Please don't hesitate to ask us questions about any doubts you might have!

Control questions

Control Question #1

If both your performance and your teammate's performance were in the Bottom 10, what is the probability you will earn 10 EUR in this experiment?

_____ Percent

Control Question #2

In the screenshot that is presented on page 3, given the probabilities that you assigned to your and your teammate's performances being in the Top 10, what is the weight that you should choose to maximize your chances of earning 10 EUR?

Control Question #3

a) Assume you and your teammate are both in the Top 10, in how many cases will you receive a Green Check from an evaluator?

_____ out of 10 cases

b) Assume you and your teammate are both in the Top 10, in how many cases will you receive a Red X from an evaluator?

_____ out of 10 cases

c) Assume exactly ONE of you or your teammate is in the Top 10, the other one is in the Bottom 10, in how many cases will you receive a Red X from an evaluator?

_____ out of 10 cases

Control Question #4

True or false? The weight you choose can also influence the type of feedback the Team Evaluators give you.

☐

True

☐

False

Control Question #5

In this experiment how are your earnings determined in Part 2? Please tick the correct answer below:

☐

One of your five weighting decisions will be selected at random for payment.

☐

All five of your weighting decisions will be selected for payment.

9.2 Wave 2: Control Treatment

Welcome to this experiment!

In this experiment you will be asked to make several decisions that determine your earnings. For showing up on time you receive 5 EUR. During the experiment you can earn additional money. How much? That depends on the decisions that you make during the experiment. Please read and follow the instructions carefully. They contain everything you need to know to participate. All decisions are taken anonymously, that is, the identity of the participants is not revealed to the other participants at any point. Also the payout is going to take place anonymously at the end of the experiment.

Please note that from now on and during the entire experiment, no communication is permitted. Throughout the experiment the use of mobile phones, smartphones, tablets or the like is prohibited. An infringement of these rules will result in exclusion from the experiment and all payments.

If you have questions at any point, please stretch your hand out of the booth and the experimenter will come to you to answer them in private.

The experiment consists of 3 parts. In the following we will explain to you Part 1 of the experiment. We will present you with the instructions for Part 2 and Part 3 after you have finished the previous parts.

Part 1 – The trivia and logic test

The computer screen will present you with a set of **15 trivia and logic questions**. Questions similar to these are often used to measure a person's general intelligence (IQ). Your task is to answer as many of these questions correctly as possible.

Your score is determined as follows. For every correct answer you receive 2.5 points, for every wrong answer you lose 1 point (you cannot end up with less than 0 points). Questions that you do not answer will not affect your score. Your score will be transformed into Euros at the end of the experiment. The exchange rate is 1 point = 0.10 EUR.

Payoff example:

Suppose you attempted 12 out of the 15 trivia and logic questions and answered 9 of them correctly and the remaining 3 incorrectly.

That means, your score is $9 \cdot 2.5 - 3 \cdot 1 = 19.5$ points

Your payments from this part of the experiment then amount to 1.95 EUR.

You will be given **10 minutes** to solve all questions. A timer in the upper right corner of your screen will inform you how much time (in seconds) is left to finish the test. You may use the backside of these instructions to take notes, if you wish to do so.

You will not learn your score until the end of the experiment.

Attention! The test will begin shortly on the screen!

Part 2 – The team task

There are no more tests for the remainder of this experiment.

In this part of the experiment you will manage a team. Two persons from this experiment will be matched randomly as teammates. They form your team. You will have the chance to earn additional 10 EUR. How likely this is will depend on the teammates' scores in the previous test and the decisions that you are going to take in this part of the experiment.

Remember, the score depends on the number of correctly answered questions (+2.5 points each) and the number of incorrectly answered questions (-1 point each).

The identity of the teammates will never be revealed to you. The only information that the computer is going to present you with is the number of questions that the teammates *attempted* to answer in the test.

Before the start of Part 2 the computer has separately compared the two teammates' scores with the *same* randomly selected group of 19 other participants in today's experiment:

- **Teammate 1** has been compared with the 19 participants and is either **ranked in the Top 10 or Bottom 10** of the 20 performances.
- Similarly, **Teammate 2** has been compared with the same 19 participants and is similarly **ranked in either the Top 10 or Bottom 10**.
-

What determines whether you earn the additional 10 EUR in this part of the experiment?

- If both Teammate 1's and Teammate 2's scores on the test were in the **Top 10**, you will earn the 10 EUR for sure (100% probability of winning).
- If both Teammate 1's and Teammate 2's scores on the test were in the **Bottom 10**, you will NOT earn the 10 EUR for sure (0% probability of winning).
- If ONE of either Teammate 1's score or Teammate 2's score was in the Top 10, and the other was in the Bottom 10, **your probability of earning the 10 EUR depends on you choosing a "weight" from 0 to 100** that you assign to Teammate 1's performance in the team.

Being in the Top 10 indicates that one is in the top 50% (top half) of performers. Being in the Bottom 10 indicates that one is in the bottom 50% (bottom half) of performers.

We will now explain this "weight" in more detail.

Part 2 – Choosing a weight

If ONE of either Teammate 1's score or Teammate 2's score was in the Top 10, and the other was in the Bottom 10, the probability of earning the 10 EUR is calculated as follows:

- If Teammate 1 scored in the Top 10 (and Teammate 2 did not): $\sqrt{\text{Weight}/100}$
- If Teammate 2 scored in the Top 10 (and Teammate 1 did not): $\sqrt{(100 - \text{Weight})/100}$.

Don't worry if this looks confusing. **You will be asked to enter the probabilities you believe Teammate 1 and Teammate 2 scored in the Top 10.** Then the computer will automatically calculate the weight on Teammate 1's performance that gives you the highest probability of earning the 10 EUR.

To guarantee the highest probability of earning the 10 EUR, you should choose a weight **closer to 100 if you believe Teammate 1 was more likely than Teammate 2 to score in the Top 10**, and **closer to 0 if you believe Teammate 2 was more likely to score in the Top 10**. If you think that both Teammate 1 and Teammate 2 have the same likelihood of scoring in the Top 10, then the optimal weight would be 50.

Here's an example of how it works. Suppose you believe that there is a 25% chance Teammate 1's performance was in the Top 10, but you think there's a 50% chance Teammate 2's performance was in the Top 10.

Enter your estimates for the probability Teammate 1's performance was in the Top 10, and for the probability Teammate 2's performance was in the Top 10.

In order to have the greatest chance of earning the 10 EUR, you should enter these probabilities as accurately as possible.

Probability Teammate 1's performance was in the Top 10:

25 %

Probability Teammate 2's performance was in the Top 10:

50 %

Click the slider to adjust the probabilities.

CONTINUE

(You will be able to go back and change your estimates)

	Teammate 2 Top 10	Teammate 2 Bottom 10
Teammate 1 Top 10	12 %	13 %
Teammate 1 Bottom 10	38 %	37 %

Given what you have entered in the slider, the computer will estimate the probability of each of the following four scenarios in the table at the bottom of the screen, by assuming that Teammate 1's and Teammate 2's scores are independent. The four scenarios are:

- 1) Teammate 1 and Teammate 2 are both in the Top 10 [top left]
- 2) Teammate 1 is in the Top 10 and Teammate 2 is in the Bottom 10 [top right]
- 3) Teammate 1 is in the Bottom 10 and Teammate 2 is in the Top 10 [bottom left]
- 4) Teammate 1 and Teammate 2 are both in the Bottom 10 [bottom right]

Note that these four scenarios cover all of the possibilities. Therefore they must sum up to 100%.

If you would like, you can adjust any of the probabilities in the table, if you think one scenario is more likely than another. For this, you simply have to click on the "+" and "-" signs that appear next to the figures in the table. You can change the probabilities in the table as many times as you can until you decide that they reflect your true belief. To be able to proceed to the next screen you will only have to make sure that the 4 probabilities sum up to 100 % before you click "CONTINUE".

Given the above probabilities, this is how the information would be presented to you on the screen:

Enter your estimates for the four probabilities (both Top 10, Teammate 1 Top 10 only, Teammate 2 Top 10 only, both Bottom 10).
In order to have the greatest chance of earning the 10 EUR, you should enter these probabilities as accurately as possible.

Probability Teammate 1's performance was in the Top 10: 25 %

Probability Teammate 2's performance was in the Top 10: 47 %

Click here to decrease the probability you believe Teammate 1 and Teammate 2 are both in the Top 10.

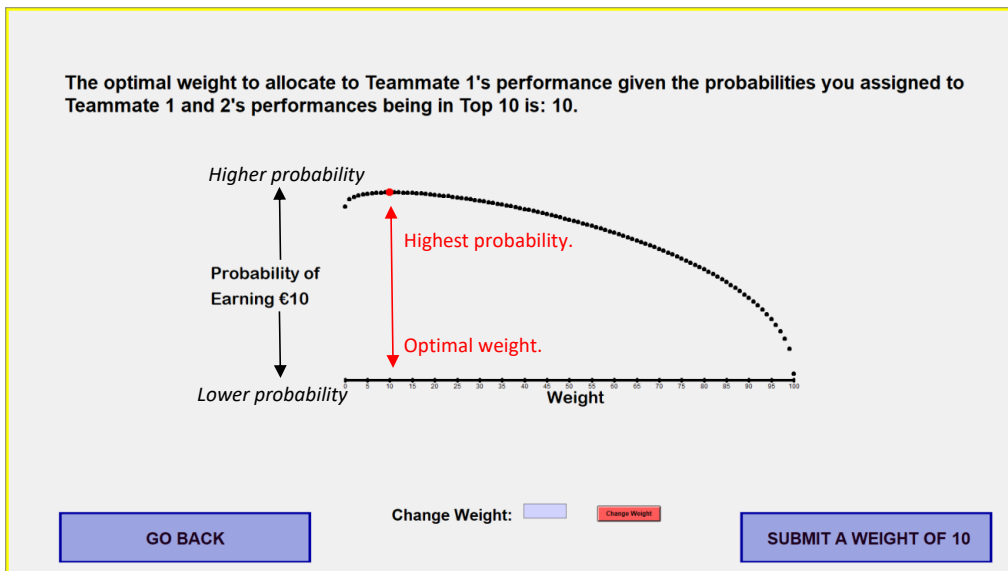
Click here to increase the probability you believe Teammate 1 is in the Top 10 but Teammate 2 is not.

	Teammate 2 Top 10	Teammate 2 Bottom 10
Teammate 1 Top 1	- 12 % +	- 13 % +
Teammate 1 Bottom 1	- 35 % +	- 37 % +

SUM: 97 %
Sum cannot be < 100%

GO BACK

Once you click "CONTINUE", the computer will calculate the probability of earning the 10 EUR for *every* possible weight from 0 to 100 and inform you which weight you should *optimally* choose given the probabilities that you enter.



Summary

In order to have the best chance of earning the 10 EUR, it pays to estimate Teammate 1's and Teammate 2's performance as accurately as possible. The computer then calculates the optimal weight for you.

If you are satisfied with this weight, you can "confirm" this weight on the following screen and submit it for evaluation.

OR

If you want to put a different weight, you may either do this directly,

OR

You may go back to the previous screen and enter different probabilities with which you expect that Teammate 1's and Teammate 2's performances were in the Top 10. Based on this, a new weight will be calculated and displayed to you.

Feedback

After you will have submitted the weight, on the next screen, a Team Evaluator will give your team either a "Green Check" or a "Red X", which is related to Teammate 1's performance and Teammate 2's performance.



However note that Team Evaluators can sometimes make mistakes:

- If both teammates scored in the Top 10, a Team Evaluator will give your team a Green Check in 9 out of 10 cases. In 1 out of 10 cases he will give your team a Red X.
- If one of the teammates scored Top 10 and the other one scored in the Bottom 10, a Team Evaluator will give your team a Green Check in 5 out of 10 cases. In other 5 out of 10 cases, he will give your team a Red X.
- If both teammates scored in the Bottom 10, a Team Evaluator will give your team a Green Check in 1 out of 10 cases. In 9 out of 10 cases he will give your team a Red X.

A **Green Check** indicates that it is **unlikely** that **both teammates** scored in the **Bottom 10**, while a **Red X** indicates that it is **unlikely** that **both teammates** scored in the **Top 10**.

That means, while a team evaluator gives you some information about Teammate 1's and Teammate 2's performances, you do not learn their performances for sure.

What happens next?

After receiving the Team Evaluator's feedback the computer will ask you once more to **enter the probabilities you believe Teammate 1 and Teammate 2 scored in the Top 10**. Then the computer will again automatically calculate the weight that gives you the highest probability of earning the 10 EUR.

As before, you may plug in and try out different numbers. Once you are satisfied with the weight that the computer calculates for you – or, alternatively, once you have entered a weight yourself – you can “confirm” this weight and submit it for evaluation.

This procedure (evaluator feedback – entering performance expectations and assigning weights) will be repeated 4 times. So that, in total, your team receives feedback from 4 different evaluators, after which the computer asks you to enter your beliefs and weights again. In each round you have the chance to earn the 10 EUR in the way described above.

Please note: Each time you receive feedback from a Team Evaluator, it is *completely independent* of the feedback your team previously received. Also note that the weight(s) you have chosen have no effect on feedback. Feedback is just about Teammate 1's and Teammate 2's performances.

What determines your earnings today?

Besides your earnings in the trivia and logic test in Part 1 of the experiment, you will be paid for Part 2 and for Part 3. **One out of the five** weighting decisions that you enter in Part 2 (your initial one and the four that you enter after your team has received evaluator feedback) will be randomly selected for payment. **Thus, every decision you make is important as it may affect your earnings if it is chosen for payment.** Part 3 will be explained after Part 2 is over, but don't worry, your decisions in Part 2 will not affect your earnings in Part 3.

At the end of the experiment the computer will inform you about its selection on your screen.

On the next page, we will ask you to answer some control questions. This is to make sure that all participants understand the instructions of this experiment. Once all participants have answered all questions correctly, Part 2 of the experiment will start.

If you have questions about any of the instructions (now or later), please stretch your hand out of the booth and the experimenter will come to you to answer them in private. Please don't hesitate to ask us questions about any doubts you might have!

Control questions

Control Question #1

If both teammates' performances were in the Bottom 10, what is the probability you will earn 10 EUR in Part 2?

_____ Percent

Control Question #2

In the lower of the two screenshots that are presented on page 4, given the probabilities that you assigned to Teammate 1's and Teammate 2's performances being in the Top 10 (in the given example: 25% and 50%, respectively), what is the weight that you should choose to maximize your chances of earning 10 EUR?

Control Question #3

a) Assume both teammates are in the Top 10, in how many cases will your team receive a Green Check from an evaluator?

_____ out of 10 cases

b) Assume both teammates are in the Top 10, in how many cases will your team receive a Red X from an evaluator?

_____ out of 10 cases

c) Assume exactly ONE of the teammates is in the Top 10, the other one is in the Bottom 10, in how many cases will your team receive a Red X from an evaluator?

_____ out of 10 cases

Control Question #4

True or false? The weight you choose can also influence the type of feedback the Team Evaluators give your team.

☐

True

☐

False

Control Question #5

In this experiment how are your earnings determined in Part 2? Please tick the correct answer below:

☐

One of your five weighting decisions will be selected at random for payment.

☐

All five of your weighting decisions will be selected for payment.

Part 3 – The new team task

In this final part of the experiment, everything is the same as in Part 2. Just like before, you will be asked to enter the probabilities that you believe Teammate 1 and Teammate 2 scored in the Top 10. Then the computer will automatically calculate the weight that gives you the highest probability of earning 10EUR.

As before, you will next receive feedback, and then once more you can enter the probabilities you believe that Teammate 1 and Teammate 2 scored in the Top 10. This procedure – entering performance expectations and assigning weights – will be repeated 5 times.

The only difference to Part 2 is that you will have an opportunity to change Teammate 2 that Teammate 1 is matched with:

Changing the teammate

If you change to a new Teammate 2, that person will be randomly selected from this session, and his or her score will be compared with the same 19 other participants as Teammate 1's score was compared to, just as before.

If you change to a new Teammate 2, the computer will present you with the number of questions that the new Teammate 2 attempted to answer in the test in Part 1.

If you do not change to a new Teammate 2 but stay with the old Teammate 2, you will be reminded of the number of questions the old Teammate 2 attempted to answer in the test in Part 1.

As before, your decisions do not affect the payoffs of any of the old or new teammates (and they will never learn about your decisions). Also, the new matching depends ONLY on your decisions, the old or new teammate's decisions are completely irrelevant.

We will now explain to you in detail how you can change Teammate 2:

Changing Teammate 2 comes at a price. On the first screen you will be asked to enter the MAXIMUM amount (from 0.00 EUR to 5.00 EUR) you are willing to pay to change Teammate 2.

The computer has already randomly generated the true price for changing teammates. It is a random amount between 0.00 EUR and 5.00 EUR, too. After entering your maximum amount, the computer will reveal the true price to you.

- If the true price is lower than your stated maximum amount, then Teammate 2 will change and you will pay this true price.
- If the true price is higher than your stated maximum amount then Teammate 2 will not change. You will not pay anything.

As you can see, the best you can do is to enter the highest amount you would be willing to pay to change Teammate 2. If you enter too high an amount, you might pay more than you wanted to change Teammate 2. If you enter too low an amount, you might not change Teammate 2 for a price you'd have been willing to pay.

There are no right or wrong answers, it just depends how much you would be willing to pay.

→ If you *wouldn't be willing to pay anything* (even 0,01 EUR) to change Teammate 2, then enter 0.

→ If you would be willing to pay up to 2 EUR, then enter 2 EUR.

→ If you would be willing to pay *any* price, then enter 5 EUR.

What determines your earnings today?

In addition to your earnings in the test in Part 1 and one randomly selected decision in Part 2, you will be paid for one randomly chosen decision in Part 3. **Every decision you make is important as it may affect your earnings if it is chosen for payment.**

At the end of the experiment the computer will inform you about its selection on your screen.

Just like Part 2, one of your weighting decisions will be selected for payment in Part 3. If this decision involved a new Teammate 2, the price for changing to that teammate will be subtracted from your Part 3 earnings

If you have questions about any of the instructions (now or later), please stretch your hand out of the booth and the experimenter will come to you to answer them in private. Please don't hesitate to ask us questions about any doubts you might have!