



Garbling an evaluation to retain an advantage[☆]

Ascensión Andina-Díaz^a , José A. García-Martínez^b

^a Dpto. Teoría e Historia Económica, Universidad de Málaga, Spain

^b Dpto. Estudios Económicos y Financieros, Universidad Miguel Hernández de Elche, Spain

ARTICLE INFO

JEL classification:

C72
D82
D83

Keywords:

Social comparisons
Expert dissent
Heterogeneous expertise
Career concerns
Probability of feedback

ABSTRACT

We study information transmission in a model of career concerns in which experts evaluate their worth based on social comparisons. There are two experts, each of whom receives an informative signal about the state of the world and makes a statement to the principal. The quality of the signal each expert receives is unknown to the other players, and the experts differ in the prior that their signal is fully informative. Accordingly, we speak of the stronger and the weaker expert, where the stronger expert is ex-ante more likely to receive a better signal. We show that expert heterogeneity and social comparisons drive expert dissent. We identify an incentive for the stronger expert to deliberately misreport an informative signal in order to sabotage the weaker expert, garble the principal's evaluation, and maintain her initial advantage. In equilibrium, this expert may even completely contradict her signal and the decision of the other expert. This result suggests a new rationale for social dissent that may help shed light on current polarization trends.

1. Introduction

It is human nature to compare ourselves to others. Sometimes unconsciously, human beings tend to evaluate our own social and personal achievements based on how we stack up against others. We do it on a daily basis and across multiple dimensions, from success and intelligence to wealth and attractiveness.

Though a regular and well-established phenomenon, our knowledge of the effects of interpersonal comparisons on certain domains of human behavior is still somewhat limited. Research in psychology has shown that people who regularly compare themselves to others may find motivation to improve, but may also experience feelings of dissatisfaction and guilt, and engage in negative behaviors such as lying (Gibbons and Buunk 1999, White et al. 2006). Experimental literature has shown effects of social comparisons on individual effort, finding higher effort levels and altruistic behavior in response to social preferences, but also deception and sabotage in competitive environments (Fehr and Schmidt 1999, Charness and Rabin 2002, Harbring et al. 2007, Edelman and Larkin 2015). With a focus on a different domain, this paper examines the value of social comparisons on the common human aspiration of being perceived as well-informed and good at one's job. More precisely, we investigate how measuring our personal worth based on how we stack up against others, affects our speech and the quality of the information we

[☆] We thank Mikhail Drugov, Gilat Levy, Melika Liporace, Miguel A. Meléndez-Jiménez, Ignacio Ortuño-Ortín, Nicola Pavoni, Raghul Venkatesh, Dimitrios Xefteris, seminar participants at the Universidad de Málaga and the audience at SING 2021, EEA-ESEM 2021, ASSET 2021, RES 2022, SAEe 2022, NICEP Conference 2023, Oligo 2023, and 2023 Workshop for Women in Political Economy for useful comments. We also thank the editor and two anonymous referees for many comments that helped improve the exposition of the model and the robustness of the results. We gratefully acknowledge the financial support from the Ministerio de Ciencia, Innovación y Universidades, Spain (MCIU/AEI/FEDER, UE) through projects PID2021-127736NB-I00 and PID2022-137211NB-I00. The usual disclaimer applies.

* Corresponding author.

E-mail addresses: aandina@uma.es (A. Andina-Díaz), jose.garciam@umh.es (J.A. García-Martínez).

<https://doi.org/10.1016/j.eurocorev.2024.104940>

Received 12 March 2024; Received in revised form 6 November 2024; Accepted 20 December 2024

Available online 3 January 2025

0014-2921/© 2025 The Authors. Published by Elsevier B.V. This is an open access article under the CC BY-NC license (<http://creativecommons.org/licenses/by-nc/4.0/>).

transmit.

There are at least three reasons why agents with private information (hereafter experts) might assess their personal worth on the basis of peer performance. First, the market's evaluation system may use some sort of social comparison. This is the case, for example, when promotion opportunities and career prospects depend on relative expertise. An example of this is Tipranks.com, where stock market analysts are ranked according to their relative performance versus one another. Another example is the political arena, where politicians of either different or the same party aspire to move up. Second, experts may compete for the attention of a listener who has a limited endowment of time. This is an increasingly relevant phenomenon in modern Western societies, where lack of time is pushing speakers more and more towards the use of strategies that attract the attention of a busy audience. An example is competition for attention in social media or any other forum. Third, rather than external pressure, experts may feel an internal voice to prove themselves better than others and self-impose this competitive spirit. This idea goes back at least to the Social Comparison Theory (Festinger, 1954), according to which individuals evaluate their own opinions and abilities by comparing themselves to others.¹

Our central finding is that social comparisons can lead to social dissent and confrontation of speeches, even in common-value contexts where experts have access to the same information and there is no preferred or popular action. We find that the probability of dissent increases in expert heterogeneity and it is at its highest when the likelihood that the audience learns the underlying state of the world is not very high. Our rationale for dissent differs from the mechanism in the “forecasting contests” literature (Ottaviani and Sørensen 2006c, Lichtendahl et al. 2013, Banerjee 2021), where experts differentiate their speeches in the hope of reducing the number of competitors with whom they share a prize; and also from the “anti-herding” literature (Effinger and Polborn 2001, Avery and Chevalier 1999, Levy 2004), where experts contradict an established opinion in order to signal ability. Common to these papers is that experts overweight their private information and “exaggerate” own signals. In contrast, we find that experts may act against private informative signals by choosing to discard them and deliberately lie. This may be optimal if experts are perceived to be heterogeneous in reputation. In this case, stronger players — who have little to gain from being proven right — may gain more by contradicting the weaker rival and sabotaging the latter's reputation. Our mechanism for dissent is thus rooted in the human drive to defend and maintain power, which is consistent with empirical findings in psychology showing that individuals with competent self-images are more likely to express dissent (Niiya et al., 2021) and behave defensively (Jordan et al., 2003).

Our results provide a novel mechanism for explaining dissent in societies, suggesting that confrontation can be a means for stronger experts to undermine weaker experts and retain “power”. This observation has direct implications for the political arena, where increased voter polarization (which makes politicians more and more different in the eyes of voters) may increase the ability of politicians to benefit from this strategy. It also has direct implications for situations where experts take positions and make recommendations in contexts such as climate change or Brexit, where it is difficult to have a counterfactual. Finally, our mechanism also helps accommodate empirical evidence from the evaluation of innovation and new ideas, which finds that disagreement is higher for more novel ideas, where uncertainty is also higher (Boudreau et al. 2016, Johnson and Proudfoot 2024), and epistemological discussions such as in Reiss (2020), who argues that disagreement is more likely in the “absence of evidentiary standards”.² In Section 4.3, we discuss some historical anecdotes that provide additional support for our mechanism.

The model we analyze has the following structure. Two experts with relevant information to communicate to a principal take positions on an issue. Each expert can be either of high ability (wise type), in which case she observes a fully informative signal of the state of the world, or of low ability (normal type), in which case she observes a noisy but informative signal. The type of an expert is the expert's private information. The other expert and the principal hold common priors on the probability that an expert is high ability and we allow the priors on the two experts to differ, i.e. we allow experts to be heterogeneous in their ex-ante worth or reputation. Accordingly, we talk about the stronger expert and the weaker expert, i.e., the expert with a higher and a lower probability of being high-ability, respectively. Who the stronger and the weaker experts are, is common knowledge. The experts observe a signal on the state and make a recommendation to the principal. Experts have career concerns and want the principal to draw favorable inferences about their personal worth or reputation. Experts however differ in whether they assess their personal worth based on how they stack up against others or not. In particular, we assume that experts of low ability compare themselves to others to a greater extent than experts of high ability. This is consistent with two fields of research. On the one hand, the research in psychology showing that the tendency to seek social comparison is negatively correlated with self-esteem and confidence (Festinger 1954, Gibbons and Buunk 1999, White et al. 2006).³ On the other hand, the research on social preferences, status competition and conspicuous consumption finding that upward social comparisons are more frequent than downward social comparisons (Frank 1985, Charness et al. 2013, Edelman and Larkin 2015).

We show that when the feedback on the state is not very high, stronger experts may have an incentive to misreport their signals and differentiate their advice from that of weaker experts.⁴ For sufficiently stronger experts, this incentive can be strong enough to

¹ Festinger (1954) argues that there is a drive for individuals to evaluate their opinions and abilities (Hypothesis I) and that when objective, non-social means are absent, individuals evaluate their opinions by comparison with the opinions of others (Hypothesis II).

² Existing arguments point to uncertainty aversion and fewer common templates against which to evaluate an idea, as the reasons for higher disagreement in the evaluation of novel ideas. Our mechanism complements these arguments by highlighting that for novel ideas, for which it is difficult to have a counterfactual, strategic motives that lead to dissent become more salient.

³ White et al. (2006) put it in p. 37, “People make social comparisons when they need both to reduce uncertainty about their abilities, performance, and other socially defined attributes, and when they need to rely on an external standard against which to judge themselves. The implication is that people who are uncertain of their self-worth, who do not have clear, internal standards, will engage in frequent social comparisons.” This is also consistent with Festinger (1954) (c.f. footnote 1).

⁴ This is the equilibrium behavior of stronger experts of normal type. In contrast, stronger expert of wise type and weaker experts (of any type) are honest in equilibrium.

induce the expert to always contradict her signal and the other expert's advice ([Proposition 4](#)). We show that this result is stronger the more different the experts' prior reputations are ([Corollary 1](#)), and that it holds for a wide family of payoff functions that include the ratio of reputations, the difference of reputations, and any monotonic transformation of these functions ([Theorem 1](#)). We also show that our result requires both social comparisons and heterogeneous experts. Indeed, when experts either do not compare themselves to others or they do but are rather homogeneous in reputation, full revelation of all experts' private information is always an equilibrium ([Propositions 2 and 3](#), respectively). Finally, we discuss how our results extend to variations of the model, such as wise-type experts receiving an imperfect signal or making social comparisons.

The rest of the paper is organized as follows. In [Section 2](#) we review the related literature, in [Section 3](#) we describe the model and the equilibrium concept, and in [Section 4](#) we present the results. [Section 4.1](#) considers the benchmark case where experts do not compare themselves to others and [Section 4.2](#) analyzes the novel case where they make social comparisons. In [Section 4.3](#), we elaborate on some of the assumptions of the model, provide some robustness analysis, and discuss some historical anecdotes. Finally, we conclude in [Section 5](#). The proofs of the results are in [Appendix A](#). The online Appendix B provides supplementary material for the proofs.

2. Related literature

Research in economics has examined the effects of social comparisons and social preferences on various dimensions of human behavior, such as effort provision in competitive environments ([Harbring et al. 2007](#), [Charness et al. 2013](#), [Edelman and Larkin 2015](#)), public good contribution ([Fehr and Schmidt 1999](#), [Fehr and Charness 2024](#)), consumption and status concerns ([Duesenberry 1949](#), [Hirsch 1976](#), [Frank 1985](#)). This research has shown that social comparisons can lead to social dilemmas with inefficient outcomes, such as overspending on “positional goods” like cars and clothes and underspending on “non-positional goods” like family time and savings ([Hirsch 1976](#), [Frank 1985](#)); or the destruction of economic value through reduced effort, deception, and sabotage ([Harbring et al. 2007](#), [Charness et al. 2013](#), [Edelman and Larkin 2015](#)). Similar to these papers, we find that social comparisons also introduce inefficiencies into society. In contrast to this research, our inefficiencies relate to information loss and strategic dissent.

The focus on information transmission relates our paper to the literature on career concerns for expertise. In line with [Ottaviani and Sørensen \(2001, 2006b,a\)](#), [Gentzkow and Shapiro \(2006\)](#), [Bourjade and Jullien \(2011\)](#), and [Andina-Díaz and García-Martínez \(2020\)](#), we focus on information transmission by more than one expert. Our contribution to this literature is to analyze for the first time the value of social comparisons when experts are heterogeneous. [Ottaviani and Sørensen \(2006b\)](#) is the paper closest to ours. The authors have a final section, [Section 7](#), where they consider experts that compare themselves to others. Unlike us, they consider homogeneous experts and a state of the world that is always revealed, conditions under which they obtain that social comparisons have no effect on experts' behavior.

The result that experts differentiate their advices in equilibrium relates our paper to the literature on forecasting contests. The seminal work by [Ottaviani and Sørensen \(2006c\)](#) showed that concerns for ranking may create an incentive for experts to differentiate their forecasts, a goal that experts could get by putting too much weight on their signals. Subsequent work by [Lichtendahl et al. \(2013\)](#) showed that the amount of differentiation increases in the number of forecasters, and more recently [Banerjee \(2021\)](#) extends the model to allow for conditional correlation in experts' signals. Like in contest games, these works assume an exogenous payoff function describing a winner-takes-all contest.⁵ This assumption has two important implications. First, it means the market commits ex-ante to a particular reward scheme — it does not necessarily use all information available to evaluate experts. Second, it introduces an incentive for *all* the experts to differentiate their advices — in an attempt to reduce the number of experts with whom to share the prize. As a result, this literature obtains differentiation with homogeneous experts. In contrast to this, in our model experts' payoffs are endogenously derived and the principal uses all the information available to him. This distinction is crucial for the results and shapes the nature of the differentiation result: from being a mechanism to soften competition for the prize (hence all experts can profit from) to be a way to take advantage of an initial asymmetry (hence experts can asymmetrically profit from). This guarantees that the mechanism underneath our results is different from the one in forecasting contests (see [Proposition 3](#)).

The result that experts differentiate their advice has also a flavor of the anti-herding literature. In this literature, an expert goes against an established opinion or “popular action” to show her ability. This can occur in simultaneous games ([Levy 2004](#), [Panova 2010](#)), where the popular action comes from the consideration of an unbalanced prior of the state of the world; or in sequential games, where the sequential structure endogenously makes the first action be the popular one ([Effinger and Polborn 2001](#), [Avery and Chevalier 1999](#)).⁶ Unlike this literature, our model assumes that all the states of the world are equally likely and competition is simultaneous. This guarantees that the mechanism in our paper does not relate to herding/anti-herding effects, as we have neither an unbalanced prior nor a first decision to contradict.

Last, we consider that the state of the world is not always revealed, which links our paper to the literature on the effects of transparency. Like us, [Canes-Wrone et al. \(2001\)](#), [Prat \(2005\)](#), [Levy \(2007\)](#), [Li and Madarász \(2008\)](#), [Fox and Van Weelden \(2012\)](#), and [Andina-Díaz and García-Martínez \(2020, 2023\)](#) consider that the principal does not always learn the state of the world. Unlike us, these work find a perverse effect of transparency. In our case, however, transparency always disciplines, as in [Gentzkow and Shapiro \(2006\)](#).

⁵ This is standard assumption in the literature on contests and tournaments. See [Lazear \(1981\)](#), [Green and Stokey \(1983\)](#), and [Nalebuff and Stiglitz \(1983\)](#).

⁶ Recently, this insight has been used by the literature on political economy that looks at the behavior of career concerned challengers that want to make the incumbent look bad ([Glazer 2007](#), [Fox and Van Weelden 2010](#), [Ashworth and Shotts 2011](#), [Buisseret 2016](#)). This literature considers sequential structures — a challenger/opposition-party plays after the incumbent/executive has taken the action — from where anti-herding motives arise.

3. The model

We consider a model between two experts $i \in \{1, 2\}$ (she) with career concerns and one principal (he). There is a binary state of the world $\omega \in \{L, R\}$ and a binary set of actions $a_i \in \{\hat{l}, \hat{r}\}$. We assume that the two states are equally likely.⁷

Experts issue recommendation (take an action) on the state of the world. Actions are simultaneous. Prior to issuing recommendation, expert i receives a private signal $s_i \in \{l, r\}$ on the state of the world. Depending on the type of expert i , $t_i \in \{W, N\}$, her signal of the state s_i is either perfectly accurate or noisy though informative. In particular, $P(s_i = \omega | t_i = W) = 1$ if the expert is wise-type (high ability) and $P(s_i = \omega | t_i = N) = \gamma \in (\frac{1}{2}, 1)$ if the expert is normal-type (low ability).⁸ Types of experts are i.i.d. and signals are i.i.d. conditional on the state. Each expert knows her type, which is her private information. The other players have a prior about the probability that an expert is a wise type. Let $\alpha_i \in (0, 1)$ be the prior that expert i is wise-type and $1 - \alpha_i$ be the prior she is normal-type. We assume $\alpha_1 \geq \alpha_2$, i.e., ex-ante expert 1 has a higher (or equal) probability of being wise-type than expert 2. Hereafter, we refer to expert 1 as the *stronger* expert and to expert 2 as the *weaker* expert.

We define the strategy of an expert as a mapping that associates with every possible type and signal of an expert a probability distribution over the space of actions. We denote by $\sigma_t^i(s)$ the probability that expert i of type t takes the action a that corresponds to her signal s . In other words, $\sigma_t^i(l) = P_t^i(\hat{l} | l)$ and $\sigma_t^i(r) = P_t^i(\hat{r} | r)$, for $i \in \{1, 2\}$ and $t \in \{W, N\}$. Then, $1 - \sigma_t^i(l) = P_t^i(\hat{r} | l)$ and $1 - \sigma_t^i(r) = P_t^i(\hat{l} | r)$ is the probability that expert i of type t takes the action a that mismatches her signal s .

After experts issue recommendation and before the principal forms a belief about the types of the two experts, there is a probability that the principal learns the state. Let $\mu > 0$ denote this probability that we refer to as the *probability of feedback*. We denote by $X \in \{L, R, \emptyset\}$ the feedback received by the principal, with $X = \emptyset$ indicating that there is no feedback and $X = L$ (R) indicating that the principal learns that the state is L (R). The principal observes the vector of actions (a_1, a_2) and feedback X and, based on this information, updates his beliefs about each of the experts' types. Let $\hat{\alpha}_i(a_i, a_j, X)$ denote the principal's posterior probability that expert i is type W , given (a_i, a_j) and X , with $i, j \in \{1, 2\}$ and $i \neq j$.

Experts have career concerns and want the principal to make favorable inferences about their personal ability and worth. Types of experts however differ in the extend to which they assess their personal worth based on how they stack up against others. In line with the Social Comparison Theory (Festinger, 1954), we assume that experts of type N assess their payoff on the basis of peer performance, whereas experts of type W may not. In particular, in the main body of the paper we assume that experts of type W do not compare themselves to others. In Section 4.3, we relax this assumption and allow wise-type experts to also care about social comparisons.⁹

The payoff of an expert i of type W that does not compare herself to others is given by $g(\hat{\alpha}_i(a_i, a_j, X))$, with $g(\hat{\alpha}_i) : [0, 1] \rightarrow \mathbb{R}_+$ being increasing in $\hat{\alpha}_i$. Note that this payoff only depends on the expert i 's posterior $\hat{\alpha}_i$, which nevertheless depends on the actions of the two experts. This is clear when $X = \emptyset$, in which case action a_j may contain information about the state; hence, it may be informative about i 's ability. In contrast to this, the payoff of an expert i of type N that compares herself to others depends on the two posteriors of the two experts and is given by $f(\hat{\alpha}_i(a_i, a_j, X), \hat{\alpha}_j(a_j, a_i, X))$, with $f(\hat{\alpha}_i, \hat{\alpha}_j) : [0, 1]^2 \rightarrow \mathbb{R}_+$ being C^1 , increasing in $\hat{\alpha}_i$ and decreasing in $\hat{\alpha}_j$.¹⁰ Note that this payoff increases in the expert's posterior and decreases in the opponent's posterior, implying that the distance away from the opponent's posterior matters. This feature suggests that with social comparisons, experts may choose to sabotage the opponent, seeking to undermine the posterior of the rival. However, the capacity to sabotage others may be asymmetric across experts.

Our equilibrium concept is Perfect Bayesian Equilibrium. We say that $(\sigma_W^i(l)^*, \sigma_W^i(r)^*; \sigma_N^i(l)^*, \sigma_N^i(r)^*)$ is an *equilibrium strategy profile* of expert $i \in \{1, 2\}$ if given the equilibrium strategy of expert j and the principal's consistent beliefs, $\sigma_t^i(l)^*$ maximizes the expected payoff of expert i of type t after observing signal l , and $\sigma_t^i(r)^*$ does it after signal r .

In the analysis that follows, we restrict attention to the use of symmetric strategies. For an expert $i \in \{1, 2\}$ of type $t \in \{W, N\}$, we say that the expert uses a *symmetric strategy* when $\sigma_t^i(l) = \sigma_t^i(r) = \sigma_t^i$. The restriction to symmetric strategies implies we analyze situations where an expert takes the action that corresponds to her signal with the same probability across the two information sets, $s = l$ and $s = r$. This is a natural restriction in our model, where the two information sets that correspond to the two possible states of the world are symmetric; namely, the two states are equally likely, the signals are equally informative across states, and the probability of feedback is fixed and invariant across actions and states.¹¹

⁷ The assumption that the prior is one-half guarantees that experts' actions are driven by their private information and not by herding motives. Then, this assumption is not a limitation but a way to control for this other lying incentives and a proof that our mechanism is different from the one in the herding literature (note also that in our model, actions are simultaneous).

⁸ We relax this assumption in Section 4.3, where we discuss how our results extend to the case that wise-type experts also receive an imperfect signal.

⁹ In this more general scenario, we show that as long as high-ability experts rely less on social comparisons than low-ability experts, our main finding holds true.

¹⁰ A function f is said to be of class C^1 in $[0, 1]^2$ if its 1-partial derivatives exist at all points of $[0, 1]^2$ and are continuous.

¹¹ Two comments here. First, the assumption of symmetric strategies excludes from the analysis the consideration of pooling strategies. The potential pooling equilibria are however of little interest and so often ignored in the literature. Even more, note that in our model, as long as the experts of type W are truthful (c.f. Propositions 1, 5, and Remark 1), there is no equilibrium where the experts of type N pool at an action — as by deviating to take the other action, an expert of type N would guarantee being perceived as a high-ability type. Second, the assumption of symmetric strategies does not rule out off-the-equilibrium path beliefs. However, the key idea here is that the wise-type experts are truthful in equilibrium. In fact, note that as long as experts of type W are truthful — which is the case in this model (c.f. Propositions 1, 5, and Remark 1) —, the two actions, $\hat{l}$ and $\hat{r}$, are always sent with positive probability in the equilibrium path.

4. Results

Our first result describes the equilibrium behavior of the experts of type W . We obtain that for any payoff function $g(\hat{a}_i)$, there is an equilibrium where the wise-type experts are honest, i.e., they truthfully reveal their signal.

Proposition 1. *For any feedback $\mu > 0$ and payoff function $g(\hat{a}_i)$, there is always an equilibrium where experts of type W use the honest strategy, i.e., $(\sigma_W^{1*}, \sigma_W^{2*}) = (1, 1)$.*

The logic is quite straightforward. When the experts evaluate their achievements in absolute terms, all that matters to have a good inference from the principal is to match the state. Since the experts of type W know that their signal is perfect, there is always an equilibrium where they follow their private information.

In the light of this result, hereafter we consider that the experts of type W use the equilibrium strategy $(\sigma_W^{1*}, \sigma_W^{2*}) = (1, 1)$, and focus our analysis on the behavior of the experts of type N . In Section 4.1, as a benchmark case, we consider the scenario where the payoff function of the experts of type N is also given by $g(\hat{a}_i)$, i.e., experts of low ability neither care about social comparisons. In Section 4.2, we move to the more interesting scenario where the experts of type N do care about how they stack up against others.

4.1. Benchmark: Absolute performance

As a benchmark case, let us suppose that the experts of type N do not compare themselves to others, but simply assess their worth and ability in absolute terms, through the principal's eyes. This corresponds to the standard case in the literature of career concerns, where the payoff function of an expert only depends on her posterior. Abusing notation, let $g(\hat{a}_i(a_i, a_j, X))$ describe the payoff of expert i of type N when she does not compare herself to others. As in the case of experts W , we assume $g(\hat{a}_i) : [0, 1] \rightarrow \mathbb{R}_+$ and g is increasing in $\hat{a}_i$. For easy reference, in the Appendix we refer to this case as the *Absolute performance* (AP) case. We obtain the following result.

Proposition 2. *For any feedback $\mu > 0$, payoff function $g(\hat{a}_i)$, and reputations $\alpha_1 \geq \alpha_2$, the equilibrium is unique and it is characterized by the two experts of type N using the honest strategy, i.e., $(\sigma_N^{1*}, \sigma_N^{2*}) = (1, 1)$.*

The result says that regardless of whether experts of type N are heterogeneous or homogeneous in their initial reputation, if they do not compare themselves to others, they will follow their signals in the unique equilibrium. The logic is similar to that applying in the previous result for the experts of type W ; the only difference being that experts N are not fully confident in their information. Despite it, since the signals are informative, the argument applies. This result extends the standard positive effect of reputation (Kreps et al., 1982) to the consideration of heterogeneous experts, showing that in this case, in equilibrium, experts also report their signals honestly.

4.2. Introducing interpersonal comparisons

This section presents the main results of the paper, which describe the behavior of the experts of type N when they care about interpersonal comparisons. As opposed to the previous case, and for easy reference, in the Appendix we refer to this case as the *Relative performance* (RP) case. The analysis of this case proceeds in two steps. In Section 4.2.1, we analyze and present results for a particular functional form of $f(\hat{a}_i, \hat{a}_j)$. In Section 4.2.2, we analyze the general case.

4.2.1. A particular payoff function

Let the following functional form describe the payoff of expert i of type N :

$$f_i(a_i, a_j, X) = \frac{\hat{a}_i(a_i, a_j, X)}{\hat{a}_i(a_i, a_j, X) + \hat{a}_j(a_j, a_i, X)}. \quad (1)$$

Note that this function is of the class $f(\hat{a}_i(a_i, a_j, X), \hat{a}_j(a_j, a_i, X))$, as it satisfies $f(\hat{a}_i, \hat{a}_j) : [0, 1]^2 \rightarrow \mathbb{R}_+$, it is C^1 , increasing in $\hat{a}_i$ and decreasing in $\hat{a}_j$. The use of this functional form presents some advantages. First, it defines a zero-sum game, as $\sum_i f_i(\cdot) = 1$. This property describes *competition for attention* games, where the listener has a finite and fixed endowment of time.¹² Second, this functional form finds support in the contest theory literature, where it corresponds to the ratio-form introduced by Tullock (1980) for the contest success function. Third, this particular functional form is highly tractable, which allows us to analyze the game in detail and obtain clear-cut predictions. Finally, it is homogeneous of degree zero, which implies that the payoff is not sensitive to changes in the unit of measure of reputations.¹³

First, we consider the case of homogeneous experts and obtain the following result.

¹² An alternative interpretation of the zero-sum game is experts that compete for a promotion, where the probability of getting promotion depends on either the reputational market share of the expert (similarly to the concept of a firm's market share) or the probability that the expert wins a lottery with $\hat{a}_i + \hat{a}_j$ tickets when she has $\hat{a}_i$ tickets.

¹³ Two notes. First, a function $f(x) : \mathbb{R}^2 \rightarrow \mathbb{R}$ is homogeneous of degree zero if $f(x) = f(\lambda x) \forall \lambda > 0$. Second, for the indeterminate case of $f(0, 0)$, we assume expression (1) takes value 0.

Proposition 3. Consider $\alpha_1 = \alpha_2$ and the payoff function $f(\hat{\alpha}_i, \hat{\alpha}_j)$ in (1). For any feedback $\mu > 0$, the equilibrium is unique and it is characterized by the two experts of type N using the honest strategy, i.e., $(\sigma_N^{1*}, \sigma_N^{2*}) = (1, 1)$.

Two notes on this result. First, it shows that when the experts of type N have the same initial reputation, whether or not they compare themselves to others, does not make a difference. For a logic to this result, note that even though here experts of type N pay attention to others, since they are identical in all dimensions they cannot take advantage of any asymmetry. In this case, an expert that aims to make a good impression to the principal cannot do better than following her informative signal, as it maximizes the probability of matching the state. Second, this result is crucial to pin down and show our contribution to the literature. In particular, it proves that the mechanism driving experts' differentiation in our paper is different from the one posited by the forecasting-contest literature (Ottaviani and Sørensen 2006c, Lichtendahl et al. 2013, Banerjee 2021). The reason is that this literature obtains expert dissent with identical experts (see our discussion in Section 2), something that can never occur in our case.

Next, we consider heterogeneous experts, i.e., $\alpha_1 > \alpha_2$. The expressions of the thresholds μ_1 , μ_2 , and $\bar{\alpha}$, and the equilibrium probability, x , are defined in the proof.¹⁴

Proposition 4. Consider $\alpha_1 > \alpha_2$ and the payoff function $f(\hat{\alpha}_i, \hat{\alpha}_j)$ in (1). There exists μ_1 and μ_2 , with $\mu_1 < \mu_2$ and $\mu_2 \in (0, 1)$, such that in the unique equilibrium:

- If $\mu > \mu_2$, then $(\sigma_N^{1*}, \sigma_N^{2*}) = (1, 1)$.
- If $\mu_1 < \mu < \mu_2$, then $(\sigma_N^{1*}, \sigma_N^{2*}) = (x, 1)$, with $x \in (0, 1)$.
- If $\mu < \mu_1$, then $(\sigma_N^{1*}, \sigma_N^{2*}) = (0, 1)$, where $\mu_1 > 0$ if and only if $\alpha_1 > \bar{\alpha}$, with $\bar{\alpha} \in (\alpha_2, 1)$.

Proposition 4 identifies different scenarios according to the probability of feedback μ . A common feature to all the scenarios is that the weaker expert of type N always follows her signal. More interesting is the behavior of the stronger expert of type N . We observe that except for the case in which the probability of feedback is sufficiently high (i.e., $\mu > \mu_2$), with interpersonal comparisons there is an incentive for the stronger expert of type N to misreport and contradict her signal, with this incentive increasing as the probability of feedback decreases. In the limit, when $\mu < \mu_1$, which requires a stronger enough expert, we obtain that this expert always contradicts her signal.

To understand how the lying incentive of the stronger expert N is sustained in equilibrium, it is useful to consider an arbitrarily high probability that the stronger expert is of type W and an arbitrarily low probability of feedback, in which case an expert's action is (most likely) judged against the other expert's action.¹⁵ In this situation, the principal will update positively about the weaker expert if both experts take the same action, and negatively otherwise. There is however little room for updating about the stronger expert. This little room makes the stronger expert of type N be better by contradicting the opponent than by matching her, as when experts take different actions the principal updates negatively about both experts, but much more negatively about the weaker due to greater uncertainty about her type. The result is that the stronger expert of type N prefers to mismatch anytime she is likely enough to be of type W and the state is unlikely to be observed by the principal. The weaker expert of type N , however, cannot do better than sticking to her signal, for fear of the stronger expert being W .

This simple intuition sheds light on the mechanism behind experts' differentiation in our model, highlighting the capacity of the stronger expert of type N to exploit her initial advantage to sabotage the weaker expert, garble the principal, and maintain her advantage. This mechanism requires three key ingredients: heterogeneity of experts, low probability of feedback, and some sort of interpersonal comparisons; as pinned down in Propositions 2–4. Although seemingly restrictive, they are very natural conditions and easy to satisfy. The next corollary presents comparative static results.

Corollary 1. Thresholds μ_2 and μ_1 satisfy $\frac{\partial \mu_2}{\partial \alpha_1} > 0$, $\frac{\partial \mu_1}{\partial \alpha_1} > 0$, and $\frac{\partial \mu_2}{\partial \gamma} > 0$. Additionally, there exist $\hat{\alpha} \in (\bar{\alpha}, 1)$ and $\hat{\gamma} \in (\frac{1}{2}, 1)$ such that if $\alpha_1 > \hat{\alpha}$ and $\gamma < \hat{\gamma}$, then $\frac{\partial \mu_1}{\partial \gamma} > 0$; otherwise, $\frac{\partial \mu_1}{\partial \gamma} < 0$.

Fig. 1 below presents a graphical description of the results of Corollary 1, where top panels display the analysis for parameter α_1 and bottom panels do it for γ . We color in blue the region where the equilibrium strategy profile of experts N is $(\sigma_N^{1*}, \sigma_N^{2*}) = (1, 1)$, referred to as the *honest* equilibrium; in orange the region where $(\sigma_N^{1*}, \sigma_N^{2*}) = (x, 1)$, with $x \in (0, 1)$, referred to as a *partial-garbling* equilibrium; and in green the region where $(\sigma_N^{1*}, \sigma_N^{2*}) = (0, 1)$, referred to as the *full-garbling* equilibrium. In all the cases, the equilibrium strategy profile of experts W is $(\sigma_W^{1*}, \sigma_W^{2*}) = (1, 1)$.

There are three relevant comments. First, an increase in α_1 increases both μ_2 and μ_1 , which means that the higher the reputation of the stronger expert is, the smaller the region of μ where the equilibrium is honest.¹⁶ Second, $\mu_2 \rightarrow 0$ when $\alpha_1 \rightarrow \alpha_2$, which means that the more similar experts are in their initial reputation, the higher the region of μ where the equilibrium is honest. In the limit,

¹⁴ The expressions of thresholds μ_1 , μ_2 , and $\bar{\alpha}$ are given by expressions (14), (15), and (16) in Appendix A.1. The equilibrium probability x is the unique solution of equation $\Delta_{r,N}^{1,RP} = 0$, where $\Delta_{r,N}^{1,RP}$ is the expected gain to expert 1 of type N from taking action $\hat{r}$ rather than $\hat{l}$ after signal r under the RP system. See expression (2) in Appendix A.1. From Lemma 1, x is also the unique solution of equation $\Delta_{l,N}^{1,RP} = 0$. The probability x is a function of the parameters of the model α_1 , α_2 , γ , and μ .

¹⁵ Quite straightforward, when $X = \emptyset$, the posterior of an expert depends on the action taken by the other expert. When $X \neq \emptyset$, this is not the case, as X is sufficient. See beliefs (4)–(7) in Online Appendix B.

¹⁶ We also obtain that $\mu_1 \rightarrow \mu_2$ when $\alpha_1 \rightarrow 1$, i.e., the higher the initial reputation of the stronger expert is, the smaller the region where there is a partial-garbling equilibrium. In the limit, the equilibrium is either honest or implies fully garbling. See the proof of Proposition 4 in Online Appendix B.

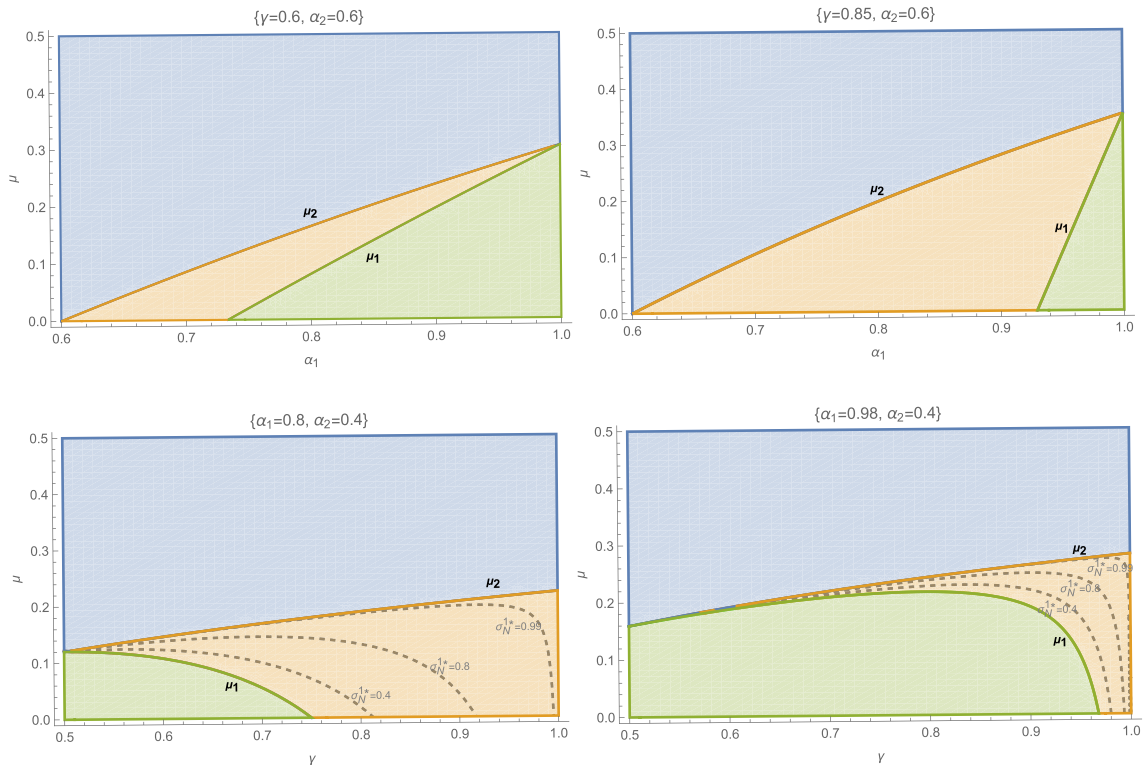


Fig. 1. Top panels represent the effect of a change in α_1 on the regions where the honest (blue), the partial-garbling (orange), and the full-garbling equilibrium (green) exist. Threshold $\bar{\alpha}$ corresponds to $\mu_1(\alpha_1) = 0$. The top-left panel considers $\alpha_2 = 0.6$ and $\gamma = 0.6$ and the top-right panel considers $\alpha_2 = 0.6$ and $\gamma = 0.85$. The bottom panels represent the effect of a change in γ on the regions where the honest (blue), the partial-garbling (orange), and the full-garbling equilibrium (green) exist. The bottom-left panel considers $\alpha_1 = 0.8$ and $\alpha_2 = 0.4$, and the bottom-right panel considers $\alpha_1 = 0.98$ and $\alpha_2 = 0.4$. The discontinuous functions represent the combinations of (γ, μ) for which the probability the stronger expert of type N follows her signal in equilibrium is $\sigma_N^{1*} \in \{0.4, 0.8, 0.99\}$. (For interpretation of the references to color in this figure legend, the reader is referred to the web version of this article.)

this is the unique equilibrium (Proposition 3).

Third and last, μ_2 increases in γ ; however for γ arbitrarily high, μ_1 decreases in γ . This result suggests that the region of μ where the stronger expert of type N is truthful decreases in γ . However, her incentive to fully contradict her signal also decreases in γ and so, in equilibrium, garbling is only partial. Indeed, in the partial-garbling-equilibrium region, we observe that the probability that the stronger expert of type N follows her signal increases in γ (discontinuous functions $\sigma_N^{1*} \in \{0.4, 0.8, 0.99\}$). The intuition is that an increase in γ gives the stronger expert of type N more accurate information about what the weaker expert will do, which makes her more likely to contradict her signal. However, the higher frequency to contradict the signal increases the payoff from sticking to it, as types W follow the signal more frequently than types N . In equilibrium, forces must compensate and the stronger expert N contradicts with probability $0 < x < 1$.¹⁷

4.2.2. General case

This section presents the results for a general payoff function $f(\hat{\alpha}_i, \hat{\alpha}_j) : [0, 1]^2 \rightarrow \mathbb{R}_+$ being C^1 , increasing in $\hat{\alpha}_i$ and decreasing in $\hat{\alpha}_j$. We focus our attention on the case of μ being arbitrarily low, for which Proposition 4 shows that full-garbling occurs in equilibrium. In this section, we show that for any payoff function $f(\hat{\alpha}_i, \hat{\alpha}_j)$ that satisfies the *sensitivity condition* that we define next, the equilibrium of the game is unique and exhibits full-garbling.

Definition 1. A function $f(x, y) : [0, 1]^2 \rightarrow \mathbb{R}_+$ that is C^1 in $[0, 1]^2$ satisfies the sensitivity condition whenever:

$$(1) \left| \frac{\partial f(x, y)}{\partial x} \right| \leq \left| \frac{\partial f(x, y)}{\partial y} \right| \text{ if } x > y,$$

¹⁷ Regarding the effect of γ on μ_1 , we observe it is not always monotonic: whereas the left-hand side panel represents a situation in which $\frac{\partial \mu_1}{\partial \gamma} < 0$ always, the right-hand side panel represents a situation in which first $\frac{\partial \mu_1}{\partial \gamma} > 0$ and then $\frac{\partial \mu_1}{\partial \gamma} < 0$. According to Corollary 1, in the left-hand side panel we have $\alpha_1 < \bar{\alpha}$ and in the right-hand side panel we have $\alpha_1 > \bar{\alpha}$. A final comment is that $\mu_1 \rightarrow \mu_2$ when $\gamma \rightarrow 1/2$, which implies that the smaller γ , the smaller the region where the partial-garbling equilibrium exists. In the limit, the stronger expert of type N either sticks or contradicts her signal. See the proof of Proposition 4 in Online Appendix B.

$$(2) \left| \frac{\partial f(x,y)}{\partial x} \right| \geq \left| \frac{\partial f(x,y)}{\partial y} \right| \text{ if } x < y.$$

Note that the generic payoff function $f(\hat{\alpha}_i, \hat{\alpha}_j)$ exhibits $\partial f(\hat{\alpha}_i, \hat{\alpha}_j)/\partial \hat{\alpha}_i > 0$ and $\partial f(\hat{\alpha}_i, \hat{\alpha}_j)/\partial \hat{\alpha}_j < 0$. To this case, the sensitivity condition requires that whenever $\hat{\alpha}_i > \hat{\alpha}_j$, the increase in the payoff of expert i from an infinitesimal increase in $\hat{\alpha}_i$ is smaller than the increase in her payoff from the same infinitesimal decrease in $\hat{\alpha}_j$. It is straightforward to show that many payoff functions, such as $f(\hat{\alpha}_i, \hat{\alpha}_j) = \frac{\hat{\alpha}_i}{\hat{\alpha}_i + \hat{\alpha}_j}$, $f(\hat{\alpha}_i, \hat{\alpha}_j) = \frac{\hat{\alpha}_i}{\hat{\alpha}_j}$, $f(\hat{\alpha}_i, \hat{\alpha}_j) = \hat{\alpha}_i - \hat{\alpha}_j$, and any monotonic transformation of these functions, satisfy the sensitivity condition.¹⁸ To any function satisfying this condition, we obtain the following result:

Theorem 1. Consider $\alpha_1 > \alpha_2$ and any payoff function $f(\hat{\alpha}_i, \hat{\alpha}_j)$ that satisfies the sensitivity condition. There exists $\bar{\mu} \in (0, 1)$ and $\bar{\alpha} \in (0, 1)$, such that in the unique equilibrium $(\sigma_N^{1*}, \sigma_N^{2*}) = (0, 1)$, for any $\mu < \bar{\mu}$ and $\alpha_1 > \bar{\alpha}$.

Theorem 1 shows that as long as the probability of feedback is arbitrarily low and the stronger expert of type N is sufficiently strong, any payoff function that satisfies the sensitivity condition has a unique equilibrium, where full-garbling is the outcome. To understand this result, it is useful to examine what the sensitivity condition imposes on experts. Regarding the stronger expert, the condition tells the expert gains more from sabotaging the weaker expert and reducing her posterior $\hat{\alpha}_2$ than from guessing the state and increasing her own posterior $\hat{\alpha}_1$.¹⁹ Regarding the weaker expert instead, the condition reads opposite: It tells the weaker expert to better follow her signal than contradict it, as she gains more from increasing her posterior $\hat{\alpha}_2$ than from reducing the posterior of the opponent $\hat{\alpha}_1$. These ideas suggest that any payoff function that satisfies the sensitivity condition exhibits certain asymmetry in the power of the players. Relevant to the results, this asymmetry produces different incentives for the players: an incentive for the stronger expert to mismatch and contradict the weaker expert (as in strategic substitutes games), and an incentive for the weaker expert to match the opponent (as in strategic complements games).

4.3. Discussion

In this section, we first elaborate on some of the assumptions of the model and explore the robustness of our results to variations in these assumptions. Finally, we present some historical anecdotes that help support the empirical relevance of our mechanism.

Connection with contest games.— We start briefly elaborating on contest games, which are described as games in which contestants make a task aiming at gaining one or more *exogenous* prizes with some probability (see Corchón and Serena (2018) for a survey). Although similar in spirit to this paper, it is easy to see there are important differences between contest games and our model. To pin down these differences, it is useful to consider a simple winner-takes-all contest in which our two experts compete for a prize. As standard in the literature of contest games, consider the prize is exogenous and it is assigned to the expert that receives the higher posterior; the other expert receives nothing. The reader may recognize this set-up as fitting into the category of relative performance systems.²⁰ Despite it, this contest game has an equilibrium in which even heterogeneous experts reveal their information. To see why, note that if $\alpha_1 > \alpha_2$ and experts play the honest strategy $(\sigma_N^{1*}, \sigma_N^{2*}) = (1, 1)$, the posterior of the stronger expert will always be above the posterior of the weaker expert, with the latter having no chances to defeat the former.²¹ Then, what makes our model different from this game? The point is that while how far away the posteriors of the two experts are is not relevant in winner-takes-all contest games, it is indeed relevant in our case. The crucial idea is that winner-takes-all payoff functions are not continuous and so do not satisfy the sensitivity condition, whereas our f does. It suggests that for sabotage and garbling to be a robust equilibrium prediction, having players that compare themselves to others is not sufficient; we need that the payoff function exhibits some sort of the sensibility condition.

Wise-type experts with imperfect signals.— Now, we discuss an extension of our model where experts of type W receive a signal that is more informative than the signal of the experts of type N but not perfectly informative. Let γ_W , with $\gamma < \gamma_W < 1$, denote the accuracy of the signal of the experts of type W . Fig. 2 in Appendix A.4 presents graphical support for the robustness of our results to this alternative scenario. Like Fig. 1, we represent the region of parameters (α_1, μ) where the honest (blue), the partial-garbling (orange), and the full-garbling equilibrium (green) exist. The left panel considers $\gamma = 0.6$, $\gamma_W = 0.9$, and $\alpha_2 = 0.6$, and the right panel considers $\gamma = 0.85$, $\gamma_W = 0.9$, and $\alpha_2 = 0.6$. In line with Fig. 1, we observe that the incentive of the stronger expert of type N to contradict the signal increases in α_1 and decreases in μ .²² This proves the robustness of our results to this new scenario. Even more, we now observe that garbling can occur for higher levels of transparency. The reason for this milder condition is that when the experts of type W receive an imperfect signal, they can be proven wrong — even if they use the honest strategy. It softens the burden for a low-ability type to contradict her signal when there is transparency, as being proven wrong is no longer

¹⁸ Note that for $f(\hat{\alpha}_i, \hat{\alpha}_j) = \hat{\alpha}_i - \hat{\alpha}_j$ and its monotonic transformations, the sensitivity condition holds with equality.

¹⁹ Lemma 11 in Appendix A shows that if $\alpha_1 > \bar{\alpha}$, then posteriors satisfy $\hat{\alpha}_1 > \hat{\alpha}_2$.

²⁰ This is right, in the sense that the probability an expert wins the prize increases in her posterior and decreases in the posterior of the opponent.

²¹ This occurs regardless of the action observed by the principal. When the two actions coincide, the advantage of the stronger expert maintains (in the sense of being above the opponent), and when they do not, the stronger expert receives higher credibility and so also retains her advantage. Note also that in this contest game there are many other strategies that are equilibrium strategies. In particular, if the prior reputations of the two experts are sufficiently different, any strategy is an equilibrium strategy. The reason is that if the experts are perceived as very different in their prior reputations, there is no way for the weaker expert to receive a higher posterior than the stronger expert, which means the weaker expert is doomed to lose the contest. These observations suggest that continuity of the payoff function is important for dissent to emerge as a robust prediction.

²² Regarding the behavior of the other players, in Appendix A.4 we provide some numerical examples showing that being honest is an equilibrium strategy for these other players.

an irrefutable signal of being low-ability. As a result, we can find dissent for higher levels of transparency. This result reinforces the robustness of our results to this new scenario, suggesting that with imperfect signals, garbling can even be more frequent and occur for higher levels of transparency.

Wise-type experts with social concerns.— Here, we discuss the robustness of our results to the case that experts of type W also assess their payoff on the basis of peer performance. Let $\lambda \in [0, 1]$ be the degree to which an expert of type W compares herself to others. Definition 2 in Appendix A describes the payoff function of experts of type W in this case: when $\lambda = 1$ type W evaluates her achievements in relative terms and when $\lambda = 0$ she does it in absolute terms.²³ Proposition 5, also in Appendix A, shows that when the value of μ is sufficiently high, in particular $\mu > \mu_2$, all experts (stronger and weaker experts of both types W and N) are always honest, regardless of how much weight wise types put on social comparisons. Below threshold μ_2 , the stronger expert of type N begins to misreport her signal whereas experts of type W still remain honest. If μ decreases further and falls below μ_1 , the stronger expert of type N will always contradict her signal, whereas experts of type W remain honest. We can show that these results hold true for any λ lower than a certain threshold $\bar{\lambda} < 1$. For $\lambda > \bar{\lambda}$, there exists an additional threshold $\mu_0 \in (0, \mu_1)$ below which being honest is no longer an equilibrium strategy for the stronger expert of type W . The reason is that when transparency is very low, the gain from sabotaging the weaker rival becomes so large that it can even induce the stronger expert of type W to contradict her perfectly informative signal. Therefore, when comparing Propositions 4 and 5, we observe that the behavior of both types W and N is identical in the two propositions if $\lambda \leq \bar{\lambda}$. When $\lambda > \bar{\lambda}$, it exists a lower bound μ_0 such that the stronger expert of type W is still honest if $\mu > \mu_0 \in (0, \mu_1)$. These findings show the robustness of our main results to the consideration of experts of type W that, to some extent, also care about social comparisons.

Empirical discussion.— Finally, we discuss some empirical anecdotes. Before doing so, we want to emphasize that, given the tendency of individuals to hide and obfuscate strategic behavior, it is somewhat difficult and daring to claim that our rationale explains specific cases of dissent that we observe in real life. However, we believe that the following anecdotes fit well within the ingredients and predictions of our model and thus provide empirical evidence consistent with our results.

The first anecdote takes us to the Spanish History, to the years before the independence of the Spanish colonies in America. It deals with the reign of Charles III (1759–1788) and the different opinions that the Count of Floridablanca and the Count of Aranda had about the need to reorganize the government of the Spanish American colonies, about forty years before the first independence movements began. The Count of Floridablanca, Charles III's powerful secretary of state, repeatedly ignored letters from the Count of Aranda, ambassador in Paris and a subordinate. The Count of Aranda — signer, on the Spanish side, of the Treaty of Paris of 1783, whereby the American Revolutionary War ended — were familiar with the American reality and anticipating the possible contagion effect that the independence of the thirteen former English colonies could have on the aspirations of the Spanish territories in America, advised on the need to restructure the government of the Spanish colonies, moving towards a system closer to a federal government. The Count of Floridablanca, King Charles III's right-hand man and “omnipresent minister”, repeatedly rejected this advice on the grounds of no threat of revolt from the Spanish colonies (Pérez, 2005). Time proved that the Count of Floridablanca was wrong.

We conclude with a second anecdote that we open with a famous quote, “All my life I've known better than to depend on the experts. How could I have been so stupid, to let them go ahead?”, pronounced by President Kennedy time after the failed Bay of Pigs invasion in April 1961. The plan, devised by the CIA, was to invade Cuba, overthrow Fidel Castro, and prevent the possible spread of communism around the globe. It was early 1961 when the new and inexperienced Kennedy administration learned of the plan. The decision had to be made then or never, as the exiles were impatient and Cuba was about to receive weapons from the communist regimes in Eastern Europe. The invasion plan turned into a complete failure and became one of the great fiascos of the Kennedy administration and the CIA. Much has been written about the reasons for the fiasco. In a report commissioned from CIA Inspector General Lyman Kirkpatrick in October 1961, the CIA Inspector was very critical of the operations group that planned the attack.²⁴ His criticism speaks of partiality and the little, if any, weight given to other assessments of the plan. In fact, during the preparation of the plan, there were some voices and reports that indicated disagreement with the CIA's plan and some last-minute changes.²⁵ However, these voices came from weaker parties and they were quickly countered by the strength and power of the CIA's Director and Deputy Director — the architects of the plan — who were “skillful bureaucratic players with easy access to the president” and with career concerns and reputations to protect (Vandenbroucke, 1984).

5. Conclusion

We consider a model of career concerns with heterogeneous experts, where experts care about social comparisons. We identify an incentive for stronger experts to contradict their signal, aiming at sabotaging weaker experts, garbling the evaluation of the

²³ In the definition we assume that $f(\hat{a}_i, \hat{a}_j)$ is given by expression (1) and prove the results for this case (see Proposition 5). Then, we argue that for λ sufficiently small, our results are more general and hold for any function $f(\hat{a}_i, \hat{a}_j)$ that satisfies the sensitivity condition. See Remark 1 in Appendix A.

²⁴ He put it, “There was failure at high levels to concentrate informed, unwavering scrutiny on the project and to apply experienced, unbiased judgment to the menacing situations that developed”. See CIA Historical Review Program, Release as Sanitized 1997, p. 35.

²⁵ One example is the memorandum of March 15 written by the Joint Chiefs of Staff (JCS), where the JCS favored the original plan to land in Trinidad over the CIA's alternative to land in beaches bordering the Bay of Pigs. The JCS also expressed cautious support for the invasion plan. As General Gray, head of the JCS, later put it, “We thought other people would think that ‘a fair chance’ would mean ‘not too good’”. A second example is the memo written by Kennedy adviser Arthur Schlesinger in February 1961, where he expressed himself at the time of the preparations to be “disturbed by the scope of the CIA's invasion plans” and later regretted not speaking out louder. For both examples, see “Kennedy and the Bay of Pigs”, Kennedy School of Government, Case Program C14-80-279, pp. 5 and 7, respectively.

principal, and retaining their advantage.²⁶ If the probability of feedback is not very high, dissent occurs in equilibrium.

The results in this paper have direct implications for settings in which speakers compete for the attention of a listener. They suggest that, far from facilitating information transmission, competition for attention can have pernicious effects, namely misreported signals and contradicting messages. Our results also have implications for the conditions under which societies can reach consensus and dissent.²⁷ They suggest that while career concerns can promote consensus in the absence of interpersonal comparisons, if players compare themselves to others, the outcome is likely to be dissent. It proposes a new rationale for confrontation and polarization in the political arena, as well as for disagreement in everyday life domains such as work and family. Existing arguments explain dissent in terms of different preferences, access to different information, or the existence of a fixed price. In contrast to these arguments, our results show that expert dissent can arise in common-value contexts (hence, no different preferences), where the price is not fixed (hence, no incentive to reduce the number of experts with whom to share the price), and where agents have the same information.

Crucial to our findings is that experts pay attention to others and that the distance away from the opponent matters. This suggests that if interpersonal comparisons arise because of the market evaluation system, the evaluator would do well to understand that sabotage and confrontation may be an equilibrium outcome. We recognize that the nature of certain activities, such as recruitment and promotion, requires competition and hence some form of interpersonal comparison. We also recognize that such systems have advantages, such as facilitating *sorting* between types — which is a desirable property when the interest lies in future rather than present information transmission. In this sense, the results in this paper should not be taken as a recommendation to avoid evaluation systems that involve interpersonal comparisons, but as an advice and warning not to neglect their drawbacks.

Finally, experts may pay attention to others not because of some external or formal institution, but simply because of the competitive spirit that drives much of human behavior in modern Western societies. In these cases, the results of this paper suggest that experts should be wary of paying attention to others when the market does not. In fact, if doing so, there is a risk of thinking that the expert benefits by contradicting the opponent, when the opposite is true.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgement

Funding for open access charge: Universidad de Málaga/CBUA.

Appendix A

The Appendix has the following structure. In [Appendix A.1](#), we describe the equilibrium conditions and define the instrumental functions that we use to compute the equilibrium. In [Appendix A.2](#), we derive the posterior beliefs and the expected payoffs that we use in the instrumental functions. In [Appendix A.3](#), we present the proofs of the results in the paper, and in [Appendix A.4](#), we present graphical robustness tests.

A.1. Part I: Equilibrium conditions

Let $z \in \{AP, RP\}$ describe the payoff function, with AP standing for the case where experts do not compare themselves to others (absolute performance case) and RP for the case experts do compare (relative performance case). We denote by $EU_t^{i,z}(a_i | s_i)$ the expected payoff to expert $i \in \{1, 2\}$ of type $t \in \{W, N\}$ when she takes action $a_i \in \{\hat{l}_i, \hat{r}_i\}$ and observes signal $s_i \in \{l_i, r_i\}$, under system $z \in \{AP, RP\}$. We denote by $\Delta_{s,i}^{i,z}$ the expected gain to expert i of type t from taking action $\hat{r}_i$ rather than $\hat{l}_i$ after signal s , under system z . Under symmetric strategies, we have:

$$\Delta_{r,t}^{i,z}(\sigma_W^1, \sigma_W^2; \sigma_N^1, \sigma_N^2) = EU_t^{i,z}(\hat{r}_i | r_i) - EU_t^{i,z}(\hat{l}_i | r_i) \quad (2)$$

$$\Delta_{l,t}^{i,z}(\sigma_W^1, \sigma_W^2; \sigma_N^1, \sigma_N^2) = EU_t^{i,z}(\hat{r}_i | l_i) - EU_t^{i,z}(\hat{l}_i | l_i) \quad (3)$$

We use functions (2) and (3) to define the equilibrium. We describe an equilibrium profile $(\sigma_W^*, \sigma_N^*) = (\sigma_W^{1*}, \sigma_W^{2*}; \sigma_N^{1*}, \sigma_N^{2*})$ as $(\sigma_t^{i*}, \sigma_{-i}^*)$, for $i \in \{1, 2\}$ and $t \in \{W, N\}$. Then, in a Perfect Bayesian equilibrium, the equilibrium strategy of player $i \in \{1, 2\}$ of type $t \in \{W, N\}$ is:

- $\sigma_t^{i*} = 0$ if, for all $\sigma_{-i}^i \in (0, 1]$, $\Delta_{r,t}^{i,z}(\sigma_t^i, \sigma_{-i}^*) < 0$ and $\Delta_{l,t}^{i,z}(\sigma_t^i, \sigma_{-i}^*) > 0$ (with weak inequalities, ≤ 0 and ≥ 0 respectively, for $\sigma_t^i = 0$).
- $\sigma_t^{i*} = 1$ if, for all $\sigma_{-i}^i \in [0, 1)$, $\Delta_{r,t}^{i,z}(\sigma_t^i, \sigma_{-i}^*) > 0$ and $\Delta_{l,t}^{i,z}(\sigma_t^i, \sigma_{-i}^*) < 0$ (with weak inequalities, ≥ 0 and ≤ 0 respectively, for $\sigma_t^i = 1$).
- $0 < \sigma_t^{i*} < 1$ if $\Delta_{r,t}^{i,z}(\sigma_t^i, \sigma_{-i}^*) = \Delta_{l,t}^{i,z}(\sigma_t^i, \sigma_{-i}^*) = 0$.

²⁶ This idea describes our mechanism for dissent, which inspires the title of the paper. The mechanism suggests that stronger experts serve of confrontation with weaker experts to garble the principal and maximize their own payoff. The result is expert dissent. Since the payoff of expert i is given by $f(\hat{\alpha}_i, \hat{\alpha}_j)$, with f being increasing in $\hat{\alpha}_i$ and decreasing in $\hat{\alpha}_j$, in equilibrium stronger players maximize their advantage.

²⁷ From an ex-ante point of view, we say there is dissent when experts take different actions with a probability higher than one-half. The probability of dissent is $2\gamma(1-\gamma) < 1/2$ in an honest equilibrium and it is $\gamma^2 + (1-\gamma)^2 > 1/2$ in the fully-garbling equilibrium.

A.2. Part II: Beliefs and expected payoffs

In this section, we derive the posterior beliefs and the expected payoffs that serve us build expressions (2) and (3). We assume that experts use symmetric strategies and experts of type W are honest — note this is consistent with the equilibrium behavior of experts W .

The posterior beliefs $\hat{a}_i(a_i, a_j, X)$ that players place on expert $i \in \{1, 2\}$ being of type W are:

$$\hat{\alpha}_i(\hat{l}_i, R) = \hat{\alpha}_i(\hat{r}_i, L) = 0, \quad (4)$$

$$\hat{\alpha}_i(\hat{l}_i, L) = \hat{\alpha}_i(\hat{r}_i, R) = \frac{\alpha_i}{\alpha_i + (1 - \alpha_i)(\gamma\sigma_N^i + (1 - \gamma)(1 - \sigma_N^i))}, \quad (5)$$

$$\begin{aligned} \hat{\alpha}_i(\hat{l}_i, \hat{l}_j, \emptyset) &= \hat{\alpha}_i(\hat{r}_i, \hat{r}_j, \emptyset) \\ &= \frac{\alpha_i}{\alpha_i + (1 - \alpha_i) \left(\gamma\sigma_N^i + (1 - \gamma)(1 - \sigma_N^i) + (\gamma(1 - \sigma_N^i) + (1 - \gamma)\sigma_N^i) \frac{(1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j)}{\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j))} \right)}, \end{aligned} \quad (6)$$

$$\begin{aligned} \hat{\alpha}_i(\hat{l}_i, \hat{r}_j, \emptyset) &= \hat{\alpha}_i(\hat{r}_i, \hat{l}_j, \emptyset) \\ &= \frac{\alpha_i}{\alpha_i + (1 - \alpha_i) \left(\gamma\sigma_N^i + (1 - \gamma)(1 - \sigma_N^i) + (\gamma(1 - \sigma_N^i) + (1 - \gamma)\sigma_N^i) \frac{\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j))}{(1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j)} \right)}, \end{aligned} \quad (7)$$

for all $i, j \in \{1, 2\}$, $i \neq j$. Note that when $X \neq \emptyset$, $\hat{\alpha}_i(a_i, a_j, X)$ does not depend on a_j .

The expected payoff $EU_t^i(a_i | s_i)$ to expert $i \in \{1, 2\}$ of type $t \in \{W, N\}$ for taking action $a_i \in \{\hat{l}_i, \hat{r}_i\}$ after signal $s_i \in \{l_i, r_i\}$, under system $z \in \{AP, RP\}$, is:

$$\begin{aligned} EU_t^{i,z}(\hat{l}_i | s_i) &= P_t(\hat{r}_j | \hat{l}_i, s_i)EU_t^{i,z}(\hat{l}_i, \hat{r}_j | s_i) + P_t(\hat{l}_j | \hat{l}_i, s_i)EU_t^{i,z}(\hat{l}_i, \hat{l}_j | s_i), \\ EU_t^{i,z}(\hat{r}_i | s_i) &= P_t(\hat{r}_j | \hat{r}_i, s_i)EU_t^{i,z}(\hat{r}_i, \hat{r}_j | s_i) + P_t(\hat{l}_j | \hat{r}_i, s_i)EU_t^{i,z}(\hat{r}_i, \hat{l}_j | s_i). \end{aligned}$$

We denote by $EU_t^{i,z}(a_i, a_j | s_i)$ the expected payoff to expert i of type t for taking action a_i after signal s_i , under system z , when the other expert takes action $a_j \in \{\hat{l}_j, \hat{r}_j\}$:

$$EU_t^{i,z}(a_i, a_j | s_i) = (1 - \mu)\Pi_t^z(a_i, a_j, \emptyset) + \mu \left(P_t(L | s_i, a_j)\Pi_t^z(a_i, a_j, L) + P_t(R | s_i, a_j)\Pi_t^z(a_i, a_j, R) \right),$$

where $\Pi_t^z(a_i, a_j, X)$, with $X \in \{L, R, \emptyset\}$, is the payoff of expert i under system z .

Finally, we compute probabilities $P_t(\omega | s_i, a_j)$ and $P_t(a_j | a_i, s_i)$. Regarding the former, we have:

$$P_t(\omega | s_i, a_j) = \frac{P_t(s_i | a_j, \omega)P(a_j | \omega)P(\omega)}{P_t(s_i | a_j, L)P(a_j | L)P(L) + P_t(s_i | a_j, R)P(a_j | R)P(R)}.$$

When $t = W$, $P_t(\omega | s_i, a_j)$ simplifies to $P_W(R | r_i, a_j) = P_W(L | l_i, a_j) = 1$ and $P_W(R | l_i, a_j) = P_W(L | r_i, a_j) = 0$, for $a_j \in \{\hat{l}_j, \hat{r}_j\}$. When $t = N$, this is:

$$\begin{aligned} P_N(L | s_i, \hat{l}_j) &= \frac{P_N(s_i | L)}{P_N(s_i | L) + P_N(s_i | R) \frac{(1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j)}{\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j))}}, \\ P_N(L | s_i, \hat{r}_j) &= \frac{P_N(s_i | L)}{P_N(s_i | L) + P_N(s_i | R) \frac{\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j))}{(1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j)}}, \\ P_N(R | s_i, \hat{r}_j) &= \frac{P_N(s_i | R)}{P_N(s_i | R) + P_N(s_i | L) \frac{(1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j)}{\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j))}}, \\ P_N(R | s_i, \hat{l}_j) &= \frac{P_N(s_i | R)}{P_N(s_i | R) + P_N(s_i | L) \frac{\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j))}{(1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j)}}, \end{aligned}$$

with $P_N(R | s_i, a_j) = 1 - P_N(L | s_i, a_j)$, for $a_j \in \{\hat{l}_j, \hat{r}_j\}$.

Regarding the latter, note that $P_t(a_j | a_i, s_i) = P_t(a_j | s_i)$, with:

$$P_t(a_j | s_i) = P_t(a_j | s_i, L)P_t(L | s_i) + P_t(a_j | s_i, R)P_t(R | s_i),$$

When $t = W$, this is:

$$\begin{aligned} P_W(\hat{r}_j | r_i) &= P_W(\hat{l}_j | l_i) = \alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j)), \\ P_W(\hat{l}_j | r_i) &= P_W(\hat{r}_j | l_i) = (1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j), \end{aligned}$$

and when $t = N$, this is:

$$\begin{aligned} P_N(\hat{l}_j | l_i) &= \left(\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j)) \right) \gamma + (1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j)(1 - \gamma), \\ P_N(\hat{r}_j | l_i) &= (1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j) \gamma + \left(\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j)) \right) (1 - \gamma), \\ P_N(\hat{l}_j | r_i) &= \left(\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j)) \right) (1 - \gamma) + (1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j) \gamma, \end{aligned} \quad (8)$$

$$P_N(\hat{r}_j | r_i) = (1 - \alpha_j)(\gamma(1 - \sigma_N^j) + (1 - \gamma)\sigma_N^j)(1 - \gamma) + \left(\alpha_j + (1 - \alpha_j)(\gamma\sigma_N^j + (1 - \gamma)(1 - \sigma_N^j)) \right) \gamma. \quad (9)$$

Finally, we provide the explicit expressions of Eqs. (2) and (3), where we use footnote 13. When expert i is of type W , we have:

$$\Delta_{r,W}^{i,z} = P_W(\hat{r}_j | r_i) \left((1 - \mu) \left(\Pi_i^z(\hat{r}_i, \hat{r}_j, \emptyset) - \Pi_i^z(\hat{l}_i, \hat{r}_j, \emptyset) \right) + \mu \Pi_i^z(\hat{r}_i, \hat{r}_j, R) \right) + P_W(\hat{l}_j | r_i) \left((1 - \mu) \left(\Pi_i^z(\hat{r}_i, \hat{l}_j, \emptyset) - \Pi_i^z(\hat{l}_i, \hat{l}_j, \emptyset) \right) + \mu \Pi_i^z(\hat{r}_i, \hat{l}_j, R) \right), \quad (10)$$

$$\Delta_{l,W}^{i,z} = P_W(\hat{r}_j | l_i) \left((1 - \mu) \left(\Pi_i^z(\hat{r}_i, \hat{r}_j, \emptyset) - \Pi_i^z(\hat{l}_i, \hat{r}_j, \emptyset) \right) - \mu \Pi_i^z(\hat{l}_i, \hat{r}_j, L) \right) + P_W(\hat{l}_j | l_i) \left((1 - \mu) \left(\Pi_i^z(\hat{r}_i, \hat{l}_j, \emptyset) - \Pi_i^z(\hat{l}_i, \hat{l}_j, \emptyset) \right) - \mu \Pi_i^z(\hat{l}_i, \hat{l}_j, L) \right), \quad (11)$$

and when she is of type N , we have:

$$\begin{aligned} \Delta_{r,N}^{i,z} &= P_N(\hat{r}_j | r_i)(1 - \mu) \left(\Pi_i^z(\hat{r}_i, \hat{r}_j, \emptyset) - \Pi_i^z(\hat{l}_i, \hat{r}_j, \emptyset) \right) + \\ &P_N(\hat{r}_j | r_i)\mu \left(P_N(R | r_i, \hat{r}_j) \left(\Pi_i^z(\hat{r}_i, \hat{r}_j, R) - \Pi_i^z(\hat{l}_i, \hat{r}_j, R) \right) + P_N(L | r_i, \hat{r}_j) \left(\Pi_i^z(\hat{r}_i, \hat{r}_j, L) - \Pi_i^z(\hat{l}_i, \hat{r}_j, L) \right) \right) + \\ &P_N(\hat{l}_j | r_i)(1 - \mu) \left(\Pi_i^z(\hat{r}_i, \hat{l}_j, \emptyset) - \Pi_i^z(\hat{l}_i, \hat{l}_j, \emptyset) \right) + \\ &P_N(\hat{l}_j | r_i)\mu \left(P_N(R | r_i, \hat{l}_j) \left(\Pi_i^z(\hat{r}_i, \hat{l}_j, R) - \Pi_i^z(\hat{l}_i, \hat{l}_j, R) \right) + P_N(L | r_i, \hat{l}_j) \left(\Pi_i^z(\hat{r}_i, \hat{l}_j, L) - \Pi_i^z(\hat{l}_i, \hat{l}_j, L) \right) \right), \end{aligned} \quad (12)$$

$$\begin{aligned} \Delta_{l,N}^{i,z} &= P_N(\hat{r}_j | l_i)(1 - \mu) \left(\Pi_i^z(\hat{r}_i, \hat{r}_j, \emptyset) - \Pi_i^z(\hat{l}_i, \hat{r}_j, \emptyset) \right) + \\ &P_N(\hat{r}_j | l_i)\mu \left(P_N(R | l_i, \hat{r}_j) \left(\Pi_i^z(\hat{r}_i, \hat{r}_j, R) - \Pi_i^z(\hat{l}_i, \hat{r}_j, R) \right) + P_N(L | l_i, \hat{r}_j) \left(\Pi_i^z(\hat{r}_i, \hat{r}_j, L) - \Pi_i^z(\hat{l}_i, \hat{r}_j, L) \right) \right) + \\ &P_N(\hat{l}_j | l_i)(1 - \mu) \left(\Pi_i^z(\hat{r}_i, \hat{l}_j, \emptyset) - \Pi_i^z(\hat{l}_i, \hat{l}_j, \emptyset) \right) + \\ &P_N(\hat{l}_j | l_i)\mu \left(P_N(R | l_i, \hat{l}_j) \left(\Pi_i^z(\hat{r}_i, \hat{l}_j, R) - \Pi_i^z(\hat{l}_i, \hat{l}_j, R) \right) + P_N(L | l_i, \hat{l}_j) \left(\Pi_i^z(\hat{r}_i, \hat{l}_j, L) - \Pi_i^z(\hat{l}_i, \hat{l}_j, L) \right) \right). \end{aligned} \quad (13)$$

A.3. Part III: Proofs

The first result is instrumental. It allows us to restrict the analysis to either one of the two information sets $s = l$ and $s = r$. We usually consider information set $s = r$.

Lemma 1. Under symmetric strategies, $\Delta_{r,t}^{i,z} = -\Delta_{l,t}^{i,z}$ for all $z \in \{AP, RP\}$, $t \in \{W, N\}$, and $i \in \{1, 2\}$.

Proof. Note that, under symmetric strategies, an expert is honest with the same probability after either signal $s_i \in \{l_i, r_i\}$. For type t , it implies $EU_t^{i,z}(\hat{l}_i | l_i) = EU_t^{i,z}(\hat{r}_i | r_i)$ and $EU_t^{i,z}(\hat{r}_i | l_i) = EU_t^{i,z}(\hat{l}_i | r_i)$. Additionally, since $\Delta_{r,t}^{i,z} = EU_t^{i,z}(\hat{r}_i | r_i) - EU_t^{i,z}(\hat{l}_i | r_i)$ and $\Delta_{l,t}^{i,z} = EU_t^{i,z}(\hat{r}_i | l_i) - EU_t^{i,z}(\hat{l}_i | l_i)$, see (2) and (3), then $\Delta_{r,t}^{i,z} = -\Delta_{l,t}^{i,z}$. ■

Proof of Proposition 1. In the proof we show that the honest strategy is a dominant strategy for experts of type W .

From Lemma 1, $\Delta_{r,W}^{i,z} = -\Delta_{l,W}^{i,z}$. Then, it suffices to prove that $\Delta_{r,W}^{i,AP} \geq 0$ when evaluated in the strategy profile $(\sigma_W^1, \sigma_W^2, \sigma_N^1, \sigma_N^2) = (1, 1, \sigma_N^1, \sigma_N^2)$, where $\sigma_N^i \in [0, 1]$ for $i = \{1, 2\}$. We use expression (10).

Expression (10) is a linear function in μ ; hence, it suffices to prove that $\Delta_{r,W}^{i,AP}(1, 1, \sigma_N^1, \sigma_N^2)|_{\mu=0} \geq 0$ and $\Delta_{r,W}^{i,AP}(1, 1, \sigma_N^1, \sigma_N^2)|_{\mu=1} \geq 0$. Calculations are intense in algebra but we can show that both conditions hold. See Online Appendix B for details. ■

Proof of Proposition 2. By Lemma 1, $\Delta_{r,N}^{i,AP} = -\Delta_{l,N}^{i,AP}$. Then, it suffices to prove $\Delta_{r,N}^{i,AP} > 0$ for any σ_N^i and σ_N^j , with $i, j \in \{1, 2\}$. Because of the strict inequality, it implies that expert i that observes signal r (l) is strictly better off (worse off) by sending $\hat{r}$ than $\hat{l}$. It proves that, in equilibrium, expert $i \in \{1, 2\}$ is honest and shows the uniqueness of the equilibrium.

To prove $\Delta_{r,N}^{i,AP} > 0$, we use the expression (12). Expression (12) is a linear function in μ ; hence, it suffices to prove $\Delta_{r,N}^{i,AP}|_{\mu=0} \geq 0$ and $\Delta_{r,N}^{i,AP}|_{\mu=1} > 0$. Calculations are intense in algebra but we can show that both conditions hold. See Online Appendix B for details. ■

Proof of Proposition 3. The proof is a limit case of the result in Proposition 4. It follows directly from the fact that when $\alpha_1 = \alpha_2$, threshold μ_2 is equal to 0. See expression (14). ■

Proof of Proposition 4. The proof consists of several steps and uses instrumental Lemmas 5–9. See Online Appendix B for these Lemmas.

First, by Lemma 1, $\Delta_{r,N}^{i,RP} = -\Delta_{l,N}^{i,RP}$. For convenience, in the proof we analyze expression $\Delta_{r,N}^{1,RP}$ for expert 1, and both $\Delta_{r,N}^{2,RP}$ and $\Delta_{l,N}^{2,RP}$ for expert 2.

Second, we provide an exhaustive list of all the equilibria configurations and the conditions for these equilibria in the next table:

	Equilibrium strategies		Equilibrium conditions	
1.	$\sigma_N^{1*} = 0$	$\sigma_N^{2*} = 0$	$\Delta_{r,N}^{1,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \leq 0$	$\Delta_{i,N}^{2,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \geq 0$
2.	$\sigma_N^{1*} = 0$	$\sigma_N^{2*} = 1$	$\Delta_{r,N}^{1,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \leq 0$	$\Delta_{i,N}^{2,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \leq 0$
3.	$\sigma_N^{1*} = 0$	$0 < \sigma_N^{2*} < 1$	$\Delta_{r,N}^{1,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \leq 0$	$\Delta_{i,N}^{2,RP}(\sigma_N^{1*}, \sigma_N^{2*}) = 0$
4.	$\sigma_N^{1*} = 1$	$\sigma_N^{2*} = 0$	$\Delta_{r,N}^{1,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \geq 0$	$\Delta_{i,N}^{2,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \geq 0$
5.	$\sigma_N^{1*} = 1$	$\sigma_N^{2*} = 1$	$\Delta_{r,N}^{1,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \geq 0$	$\Delta_{i,N}^{2,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \leq 0$
6.	$\sigma_N^{1*} = 1$	$0 < \sigma_N^{2*} < 1$	$\Delta_{r,N}^{1,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \geq 0$	$\Delta_{i,N}^{2,RP}(\sigma_N^{1*}, \sigma_N^{2*}) = 0$
7.	$0 < \sigma_N^{1*} < 1$	$\sigma_N^{2*} = 0$	$\Delta_{r,N}^{1,RP}(\sigma_N^{1*}, \sigma_N^{2*}) = 0$	$\Delta_{i,N}^{2,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \geq 0$
8.	$0 < \sigma_N^{1*} < 1$	$\sigma_N^{2*} = 1$	$\Delta_{r,N}^{1,RP}(\sigma_N^{1*}, \sigma_N^{2*}) = 0$	$\Delta_{i,N}^{2,RP}(\sigma_N^{1*}, \sigma_N^{2*}) \leq 0$
9.	$0 < \sigma_N^{1*} < 1$	$0 < \sigma_N^{2*} < 1$	$\Delta_{r,N}^{1,RP}(\sigma_N^{1*}, \sigma_N^{2*}) = 0$	$\Delta_{i,N}^{2,RP}(\sigma_N^{1*}, \sigma_N^{2*}) = 0$

Third, we use Lemmas 5–6 to prove that none of the following configurations: 1, 3, 4, 6, 7, and 9, are possible. Lemma 5 shows that $\sigma_N^{1*} = 1$ and $\sigma_N^{2*} < 1$ cannot occur simultaneously; as $\Delta_{r,N}^{2,RP} > 0$ for all $\sigma_N^1 \in [0, 1]$. This result rules out configurations 4 and 6. Lemma 6 shows $\Delta_{r,N}^{1,RP} > \Delta_{i,N}^{2,RP}$. This result rules out configurations 1, 3, 7, and 9. Then, in equilibrium, only configurations 2, 5, and 8 can hold. It implies expert 2 is always honest in equilibrium, i.e., $\sigma_N^{2*} = 1$. Hereafter, in the proof we consider $\sigma_N^{2*} = 1$.

Fourth, we analyze the behavior of expert 1. The analysis proceeds as follows:

(i) We show that $\sigma_N^{1*} = 1$ is the unique equilibrium strategy of expert 1 if and only if $\mu > \mu_2$. To prove this result, we use Lemmas 7–8. Lemma 7 shows that the more likely the normal-type follows signal r , the smaller the gain from taking action $\hat{r}$ rather than $\hat{l}$, i.e., $\frac{\partial \Delta_{r,N}^{1,RP}}{\partial \sigma_N^1} < 0$. Additionally, Lemma 8 shows that $\Delta_{r,N}^{1,RP}|_{\sigma_N^1=1} > 0$ if and only if $\mu > \mu_2$, with

$$\mu_2 = \frac{\alpha_1 \alpha_2 (\alpha_1 - \alpha_2) (2\alpha_2 (\gamma - 1) - 2\gamma + 1)}{(1 - \alpha_2) (\alpha_1^2 (2\alpha_2 - 1) (\gamma - 1)^2 - \gamma) + \alpha_1 \alpha_2 ((\gamma - 2)\gamma - 3\alpha_2 (\gamma - 1)^2) + \alpha_2^2 (\gamma - 1) \gamma}. \quad (14)$$

Since $\Delta_{r,N}^{1,RP}$ is positive at $\sigma_N^1 = 1$ (when $\mu > \mu_2$) and decreasing in σ_N^1 , it means $\Delta_{r,N}^{1,RP}$ is always positive in this case. In other words, we have that $\Delta_{r,N}^{1,RP} > 0$ for all $\sigma_N^1 \in [0, 1]$ if and only if $\mu > \mu_2$; hence $\sigma_N^{1*} = 1$ if and only if $\mu > \mu_2$.

(ii) We analyze the case $\mu < \mu_2$. We use Lemma 9, which consists of two points. The first point shows that $\Delta_{r,N}^{1,RP}|_{\sigma_N^1=0} < 0 \iff \alpha_1 > \bar{\alpha}$ and $\mu < \mu_1$, with

$$\bar{\alpha} = \frac{\alpha_2 (1 - \gamma) + 2\gamma - 1}{\gamma}, \quad (15)$$

and

$$\mu_1 = \frac{\alpha_1 \alpha_2 (2\alpha_2 (\gamma - 1) - 2\gamma + 1) (\gamma (\alpha_1 + \alpha_2 - 2) - \alpha_2 + 1)}{(1 - \alpha_2) (\gamma - 1) (\alpha_1^2 (\gamma - \alpha_2 (\gamma + \alpha_2 + \gamma + 1) - 1)) + \alpha_1 \alpha_2 (\alpha_2 ((\gamma - 1)\gamma + 1) + \gamma^2 + \gamma) - \alpha_2^2 (\gamma - 1) \gamma}, \quad (16)$$

with $\mu_1 < \mu_2$ (see Online Appendix B).

Since, by Lemma 7, $\Delta_{r,N}^{1,RP}$ is decreasing in σ_N^1 , the first point of Lemma 9 implies that if $\alpha_1 > \bar{\alpha}$ and $\mu < \mu_1$, then $\Delta_{r,N}^{1,RP} < 0$ for all $\sigma_N^1 \in [0, 1]$. Hence, the unique equilibrium strategy of expert 1 is $\sigma_N^{1*} = 0$ in this case.

The second point of Lemma 9 shows that $\Delta_{r,N}^{1,RP}|_{\sigma_N^1=0} > 0$ if and only if either $\alpha_1 < \bar{\alpha}$ or both $\alpha_1 > \bar{\alpha}$ and $\mu > \mu_1$ hold. It implies that if $\mu_1 < \mu < \mu_2$, then the decreasing function $\Delta_{r,N}^{1,RP}$ is positive at $\sigma_N^1 = 0$ and negative at $\sigma_N^1 = 1$, which implies $\Delta_{r,N}^{1,RP} = 0$ has a unique solution $x \in (0, 1)$. This probability x defines the unique equilibrium strategy, i.e., $\sigma_N^{1*} = x$. ■

Proof of Corollary 1. It follows from the sign of the derivatives of expressions (14) and (16). See Online Appendix B for the details. ■

Proof of Theorem 1.

By Lemma 1, $\Delta_{r,N}^{i,RP} = -\Delta_{i,N}^{i,RP}$. Hence, hereafter we focus on $\Delta_{r,N}^{i,RP}$, for $i, j \in \{1, 2\}$.

To prove the result, first we consider $\mu = 0$ and show that if α_1 is greater than a certain threshold $\bar{\alpha} < 1$, then $\Delta_{r,N}^{2,RP} > 0$ and $\Delta_{r,N}^{1,RP} < 0$. We analyze expression (12) in Online Appendix B, which corresponds to the detailed version of expression (2), and obtain the following two results:

- (1) $\Delta_{r,N}^{2,RP} > 0$ for any $(\sigma_N^1, \sigma_N^2) \in [0, 1]^2$, if $\alpha_1 > \bar{\alpha}$ and $\mu = 0$.
- (2) $\Delta_{r,N}^{1,RP} < 0$ for any $\sigma_N^1 \in [0, 1]$, if $\sigma_N^{2*} = 1$, $\alpha_1 > \bar{\alpha}$ and $\mu = 0$.

Proposition 6 in Online Appendix B proves the first result, which characterizes the conditions under which the weak expert is always honest, i.e., $\sigma_N^{2*} = 1$. To prove this result, we simplify function $\Delta_{r,N}^{2,RP}$, given by expression (12), and obtain

$$\Delta_{r,N}^{2,RP} > 0 \iff f(\hat{\alpha}_2(\hat{r}_2, \hat{r}_1, \emptyset), \hat{\alpha}_1(\hat{r}_1, \hat{r}_2, \emptyset)) - f(\hat{\alpha}_2(\hat{l}_2, \hat{r}_1, \emptyset), \hat{\alpha}_1(\hat{r}_1, \hat{l}_2, \emptyset)) > 0.$$

Then, we apply the sensitivity condition and the directional derivative of the function f to show that function f is decreasing from the point $(\hat{\alpha}_2(\hat{r}_2, \hat{r}_1, \emptyset), \hat{\alpha}_1(\hat{r}_1, \hat{r}_2, \emptyset))$ to the point $(\hat{\alpha}_2(\hat{l}_2, \hat{r}_1, \emptyset), \hat{\alpha}_1(\hat{r}_1, \hat{l}_2, \emptyset))$. It implies $\Delta_{r,N}^{2,RP} > 0$.

Proposition 7 in Online Appendix B proves the second result. It characterizes the conditions under which the strong expert always contradicts her signal, i.e., $\sigma_N^1 = 0$. The proof of this result is analogous to the proof of Proposition 6.

Finally, we apply a continuity argument and argue that by the continuity of $\Delta_{r,N}^{i,RP}$ in μ (see expression (12) in Online Appendix B), there exists $\tilde{\mu} > 0$ such that $\Delta_{r,N}^{2,RP} > 0$ and $\Delta_{r,N}^{1,RP} < 0$ for any $\mu < \tilde{\mu}$ and $\alpha_1 > \tilde{\alpha}$. This completes the proof. ■

Definition 2. Let the payoff function of experts of type W be the linear convex combination of the Absolute performance (AP) and the Relative performance (RP) payoff functions, with $\lambda \in [0, 1]$:

$$\lambda f(\hat{\alpha}_i(a_i, a_j, X), \hat{\alpha}_j(a_i, a_j, X)) + (1 - \lambda)g(\hat{\alpha}_i(a_i, a_j, X)), \quad (17)$$

with $f(\hat{\alpha}_i(a_i, a_j, X), \hat{\alpha}_j(a_i, a_j, X))$ given by expression (1). We refer to this payoff function as the *Linear combination* (LC) case.

Proposition 5. Consider $\alpha_1 > \alpha_2$, experts of type N with payoff function (1), and experts of type W with payoff function (17). There exists $\bar{\lambda} \in (0, 1)$, μ_0 , μ_1 , and μ_2 with $\mu_0 < \mu_1 < \mu_2$ and $\mu_2 \in (0, 1)$, such that for all λ :

1. If $\mu > \mu_2$, then $(\sigma_W^1, \sigma_W^{2*}, \sigma_N^1, \sigma_N^{2*}) = (1, 1; 1, 1)$.
2. If $\mu_1 < \mu < \mu_2$, then $(\sigma_W^1, \sigma_W^{2*}, \sigma_N^1, \sigma_N^{2*}) = (1, 1; x, 1)$, with $x \in (0, 1)$.
3. If $\mu_0 < \mu < \mu_1$, then $(\sigma_W^1, \sigma_W^{2*}, \sigma_N^1, \sigma_N^{2*}) = (1, 1; 0, 1)$, where $\mu_1 > 0$ if and only if $\alpha_1 > \bar{\alpha}$, with $\bar{\alpha} \in (\alpha_2, 1)$.

Additionally, $\mu_0 < 0$ if and only if $\lambda \leq \bar{\lambda}$, in which case points 1, 2, and 3 above fully characterize the equilibrium behavior. When $\lambda > \bar{\lambda}$, $\mu_0 > 0$ instead. In this case, for $\mu < \mu_0$, $(\sigma_W^1, \sigma_W^{2*}) = (1, 1)$ is not an equilibrium strategy of type W .

Proof. Proposition 4 describes the equilibrium behavior of experts of type N when experts of type W are honest. This result proves the equilibrium behavior of experts of type N in Proposition 5. Thus, the analysis that follows solely focuses on the behavior of experts of type W . First, we introduce and prove four instrumental results: Lemmas 2, 3 and 4, and Corollary 2.

Lemma 2. For an expert $i \in \{1, 2\}$ of type W with payoff function (17), $\Delta_{s,W}^{i,LC} = \lambda(\Delta_{s,W}^{i,RP}) - (1 - \lambda)(\Delta_{s,W}^{i,AP})$ with $s \in \{r, l\}$.

Proof. The result follows from expressions (10) and (11), assuming $z = LC$. ♦

Lemma 3. Under symmetric strategies, $\Delta_{r,W}^{i,LC} = -\Delta_{l,W}^{i,LC}$ for $i \in \{1, 2\}$.

Proof. It follows from Lemmas 1 and 2. ♦

Lemma 4. In the RP case, for all strategy profiles such that $\sigma_N^1 \in [0, 1]$ and $(\sigma_N^2, \sigma_W^1, \sigma_W^2) = (1, 1, 1)$, there are two thresholds μ_S and μ_W that depend on σ_N^1 such that:

1. $\Delta_{r,N}^{1,RP} \geq 0 \iff \mu \geq \mu_S$ and $\Delta_{r,W}^{1,RP} \geq 0 \iff \mu \geq \mu_S$,
 2. $\Delta_{r,N}^{2,RP} \geq 0 \iff \mu \geq \mu_W$ and $\Delta_{r,W}^{2,RP} \geq 0 \iff \mu \geq \mu_W$.
- In addition, if $\sigma_N^1 = 0$, then $\mu_S = \mu_1$, with μ_1 given by expression (16).

Proof. First, we calculate $\Delta_{r,N}^{1,RP}$ evaluated in $(\sigma_N^2, \sigma_W^1, \sigma_W^2) = (1, 1, 1)$, and obtain a linear expression in μ that depends on σ_N^1 , α_1 , α_2 , γ , and μ . For convenience, we rewrite this expression as $\Delta_{r,N}^{1,RP} = f_1 + f_2\mu$, where f_1 and f_2 are functions. Note that $\Delta_{r,N}^{1,RP} = 0 \iff \mu = \frac{-f_1}{f_2}$. Analogously, we proceed for $\Delta_{r,W}^{1,RP}$ and rewrite the expression as $\Delta_{r,W}^{1,RP} = g_1 + g_2\mu$, where g_1 and g_2 are also functions. Note also that $\Delta_{r,W}^{1,RP} = 0 \iff \mu = \frac{-g_1}{g_2}$.

It can be shown that $\frac{f_1}{f_2} = \frac{g_1}{g_2}$, despite $f_1 \neq g_1$ and $f_2 \neq g_2$. This implies there is a unique value of μ , denoted by μ_S , that makes both $\Delta_{r,N}^{1,RP} = 0$ and $\Delta_{r,W}^{1,RP} = 0$. The explicit expression of μ_S is too long to be included but it is available from the authors upon request.

Second, we show that $\Delta_{r,N}^{1,RP} > 0 \iff \mu > \mu_S$. For this, it suffices to prove that $\Delta_{r,N}^{1,RP}$ is increasing in μ . Note that $\Delta_{r,N}^{1,RP}$ is linear in μ with $\Delta_{r,N}^{1,RP}|_{\mu=0} < \Delta_{r,N}^{1,RP}|_{\mu=1}$, which proves the result. Analogously, it can be shown that $\Delta_{r,W}^{1,RP}|_{\mu=0} < \Delta_{r,W}^{1,RP}|_{\mu=1}$, which proves that $\Delta_{r,W}^{1,RP} > 0 \iff \mu > \mu_S$.

The proof of point 2. is analogous and we omit it. From here, we obtain threshold μ_W . Again, the explicit expression of μ_W is too long to be included and it is available from the authors upon request.

Finally, it is straightforward to show that $\mu_S|_{\sigma_N^1=0} = \mu_1$, with μ_1 being defined by expression (16) in Appendix A. ♦

Corollary 2. $\Delta_{r,N}^{1,RP} \geq 0 \iff \Delta_{r,W}^{1,RP} \geq 0$ and $\Delta_{r,N}^{2,RP} \geq 0 \iff \Delta_{r,W}^{2,RP} \geq 0$.

Proof. The proof follows directly from Lemma 4. ♦

We are now in position to prove the equilibrium behavior of types W in Proposition 5. We note that according to Lemma 3, $\Delta_{r,W}^{i,LC} = -\Delta_{r,W}^{i,LC}$. Then, it is sufficient to analyze $\Delta_{r,W}^{i,LC}$.

We start showing that $\Delta_{r,W}^{1,RP} \geq 0$ is a sufficient condition to establish that type W is honest in the equilibrium of Proposition 5, Points 1, 2, and 3. To show it, note that according to Lemma 2, $\Delta_{r,W}^{i,LC} = \lambda(\Delta_{r,W}^{i,RP}) - (1 - \lambda)(\Delta_{r,W}^{i,AP})$ with $i \in \{1, 2\}$. Furthermore, note that $\Delta_{r,W}^{i,AP} \big|_{\sigma_W^1, \sigma_W^2, \sigma_N^1, \sigma_N^2=1} \geq 0$ always, as $(\sigma_W^1, \sigma_W^2) = (1, 1)$ is always an equilibrium strategy in AP case (see Proposition 1). Hence, if $\Delta_{r,W}^{1,RP} \geq 0$, then $\Delta_{r,W}^{i,LC} \geq 0$ for any $\lambda \in [0, 1]$, so $(\sigma_W^1, \sigma_W^2) = (1, 1)$ in the LC case. Hence, we only need to prove that $\Delta_{r,W}^{1,RP} \geq 0$ in points 1, 2, and 3 of Proposition 5. To prove this, we use Corollary 2.

Let us first consider point 1. Regarding the normal type, we know that $\Delta_{r,N}^{1,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=1, \sigma_N^2=1} \geq 0$ and $\Delta_{r,N}^{2,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=1, \sigma_N^2=1} \geq 0$, since $\sigma_N^1 = 1$ and $\sigma_N^2 = 1$ are equilibrium strategies according to Proposition 4. Therefore, by Corollary 2, we have $\Delta_{r,W}^{1,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=1, \sigma_N^2=1} \geq 0$ and $\Delta_{r,W}^{2,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=1, \sigma_N^2=1} \geq 0$. Consequently, $(\sigma_W^1, \sigma_W^2) = (1, 1)$ is an equilibrium strategy of experts of type W .

Consider now point 2. Regarding the normal type, we have $\Delta_{r,N}^{1,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=x, \sigma_N^2=1} = 0$ and $\Delta_{r,N}^{2,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=x, \sigma_N^2=1} \geq 0$, since $\sigma_N^1 = x \in (0, 1)$ and $\sigma_N^2 = 1$ are equilibrium strategies according to Proposition 4. Hence, by Corollary 2, $\Delta_{r,W}^{1,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=x, \sigma_N^2=1} = 0$ and $\Delta_{r,W}^{2,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=x, \sigma_N^2=1} \geq 0$. It implies $(\sigma_W^1, \sigma_W^2) = (1, 1)$ is an equilibrium strategy of experts of type W .

To prove point 3, we must show that $\Delta_{r,W}^{i,LC} = \lambda(\Delta_{r,W}^{i,RP}) - (1 - \lambda)(\Delta_{r,W}^{i,AP}) \geq 0$ in the profile $(\sigma_W^1, \sigma_W^2; \sigma_N^1, \sigma_N^2) = (1, 1; 0, 1)$ when $\mu \in (\mu_0, \mu_1)$. We first consider the AP case and then the RP case.

In the AP case, it can be shown that $\Delta_{r,W}^{i,AP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=0, \sigma_N^2=1} > 0$. To prove this, we only need to show that $\Delta_{r,W}^{i,AP} \big|_{\mu=0} > 0$ and $\Delta_{r,W}^{i,AP} \big|_{\mu=1} > 0$, as $\Delta_{r,W}^{i,AP}$ is linear in μ . From Proposition 1 in Online Appendix B, making $\sigma_N^1 = 0$ and $\sigma_N^2 = 1$, we obtain $\Delta_{r,W}^{i,AP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=0, \sigma_N^2=1; \mu=0} > 0$ and $\Delta_{r,W}^{i,AP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=0, \sigma_N^2=1; \mu=1} > 0$.

Regarding the RP case, let us first consider the weaker experts ($i = 2$). We can show that if $\mu \leq \mu_1$, then $\Delta_{r,W}^{2,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=0, \sigma_N^2=1} \geq 0$. To demonstrate it, note that for the weaker expert of type N , $\Delta_{r,N}^{2,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=0, \sigma_N^2=1} \geq 0$, as $\sigma_N^2 = 1$ is an equilibrium strategy for her. Thus, by Corollary 2, $\Delta_{r,W}^{2,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=0, \sigma_N^2=1} \geq 0$, which implies $\sigma_W^2 = 1$ is an equilibrium strategy for the weak expert of type W .

Now, we focus on the stronger experts ($i = 1$). By Lemma 4, if $\mu \leq \mu_1$, then $\Delta_{r,W}^{1,RP} \big|_{\sigma_W^1=1, \sigma_W^2=1, \sigma_N^1=0, \sigma_N^2=1} \leq 0$; with equality when $\mu = \mu_1$. Therefore, when $\mu = \mu_1$ and for any $\lambda \in [0, 1]$, it holds that $\Delta_{r,W}^{1,LC} = \lambda(\Delta_{r,W}^{1,RP}) - (1 - \lambda)(\Delta_{r,W}^{1,AP}) > 0$ in $(\sigma_W^1, \sigma_W^2; \sigma_N^1, \sigma_N^2) = (1, 1; 0, 1)$, as $\Delta_{r,W}^{1,RP} = 0$ and $\Delta_{r,W}^{1,AP} > 0$. Furthermore, due to the linearity and continuity of $\Delta_{r,W}^{1,LC}$ in μ , there must exist $\mu_0 < \mu_1$ such that $\Delta_{r,W}^{1,LC} = \lambda(\Delta_{r,W}^{1,RP}) - (1 - \lambda)(\Delta_{r,W}^{1,AP}) \geq 0$ for any $\mu \in [\mu_0, \mu_1]$. Hence, $\sigma_W^1 = 1$ is an equilibrium strategy for the stronger expert of type W in this case. Consequently, if $\mu_0 > 0$, then $\Delta_{r,W}^{1,LC}$ is lower than zero for $\mu \in [0, \mu_0]$ and $\sigma_W^1 = 1$ is not an equilibrium strategy for the stronger expert of type W .

Finally, note that due to the linearity and continuity of $\Delta_{r,W}^{1,LC}$ in λ , there exists $\bar{\lambda} < 1$ such that $\Delta_{r,W}^{1,LC} = \lambda(\Delta_{r,W}^{1,RP}) + (1 - \lambda)(\Delta_{r,W}^{1,AP}) \geq 0$ for any $\lambda \in [0, \bar{\lambda}]$, because as previously stated, $\Delta_{r,W}^{1,AP} > 0$ and $\Delta_{r,W}^{1,RP} \leq 0$ if $\mu < \mu_1$. ■

Remark 1. From Definition 2, let us consider that the payoff function of experts of type W is given by (17), where $f(\hat{a}_i(a_i, a_j, X), \hat{a}_j(a_i, a_j, X))$ is any payoff function that satisfies the sensitivity condition. From Proposition 1, we know that for any payoff function $g(\hat{a}_i(a_i, a_j, X))$, there is always an equilibrium where type W uses the honest strategy, i.e., $(\sigma_W^1, \sigma_W^2) = (1, 1)$. By the linearity and continuity of expression (17) in λ , there exists $\bar{\lambda}' \in (0, 1)$ such that for any $\lambda < \bar{\lambda}'$, the honest strategy is also an equilibrium strategy of the wise type.

A.4. Part IV: Wise-type experts with imperfect signals

Here we consider the situation where experts of type W receive a signal that is more informative than the signal of experts of type N , but not perfectly informative. Let γ_W , with $\gamma < \gamma_W < 1$, denote the accuracy of the signal of a wise type expert. In Fig. 2 below, we provide graphical support for the robustness of our results to this scenario. We consider $\gamma = 0.6$ and $\alpha_2 = 0.6$ in the left panel, and $\gamma = 0.85$ and $\alpha_2 = 0.6$ in the right panel (as in Fig. 1). Additionally we set $\gamma_W = 0.9$.²⁸

²⁸ The graphical analysis is based on the new beliefs and expected payoffs that correspond to the current scenario, where experts of type W receive a signal of quality $\gamma_W \in (\gamma, 1)$. We omit these derivations for space reasons, but they are available from the authors upon request.

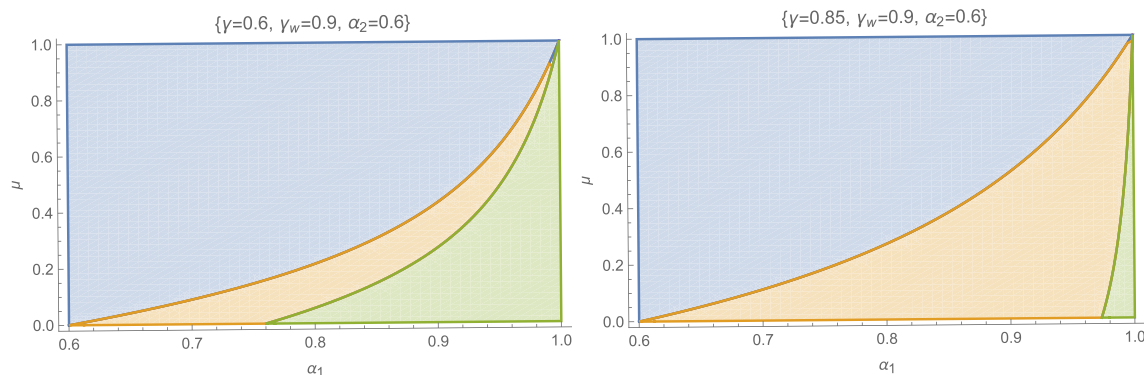


Fig. 2. We represent the regions where the honest (blue), the partial-garbling (orange), and the full-garbling equilibrium (green) exist. We consider $\gamma_W = 0.9$ and the same parameter values for γ and α_2 than we consider in the top panels of Fig. 1: $\gamma = 0.6$ and $\alpha_2 = 0.6$ in the left panel, and $\gamma = 0.85$ and $\alpha_2 = 0.6$ in the right panel. (For interpretation of the references to color in this figure legend, the reader is referred to the web version of this article.)

We represent the region of parameters (α_1, μ) where the honest (blue), the partial-garbling (orange), and the full-garbling equilibrium (green) exist. For each of these regions, we can find numerical examples showing that the honest strategy is an equilibrium strategy for all players except for the stronger expert of normal type.²⁹ Next, we present three examples that refer to the left panel, where an imaginary line is drawn at $\mu = 0.3$. Without loss of generality, we describe the equilibrium behavior of the experts after signal r .

Example 1. Let $\gamma = 0.6$, $\gamma_W = 0.9$, and $\alpha_2 = 0.6$, as in the left panel of Fig. 2. For $\mu = 0.3$, let $\alpha_1 = 0.7$, which locates the coordinates in the honest region (colored in blue). In this case, in equilibrium, $\sigma_N^{1*} = 1$. Substituting the conjectured strategy profile $(\sigma_W^1, \sigma_W^2; \sigma_N^1, \sigma_N^2) = (1, 1; 1, 1)$ into expressions (2)–(3) in Appendix A.1, we obtain $\Delta_{r,N}^{1,RP} = 0.008 > 0$, $\Delta_{r,N}^{2,RP} = 0.016 > 0$, $\Delta_{r,W}^{1,AP} = 0.149 > 0$, and $\Delta_{r,W}^{2,AP} = 0.171 > 0$; which corresponds indeed with an honest equilibrium.

Example 2. Let $\gamma = 0.6$, $\gamma_W = 0.9$, and $\alpha_2 = 0.6$, as in the left panel of Fig. 2. For $\mu = 0.3$, let us pick $\alpha_1 = 0.9$, which locates the coordinates in the partial-garbling region (colored in orange). In this case, in equilibrium, $\sigma_N^{1*} \simeq 0.251$. For the new conjectured strategy profile $(\sigma_W^1, \sigma_W^2; \sigma_N^1, \sigma_N^2) = (1, 1; 0.251, 1)$, proceeding as before, we obtain $\Delta_{r,N}^{1,RP} = 0$, $\Delta_{r,N}^{2,RP} = 0.02 > 0$, $\Delta_{r,W}^{1,AP} = 0.112 > 0$, and $\Delta_{r,W}^{2,AP} = 0.197 > 0$; which corresponds indeed with a partial-garbling equilibrium.

Example 3. Let $\gamma = 0.6$, $\gamma_W = 0.9$, and $\alpha_2 = 0.6$, as in the left panel of Fig. 2. For $\mu = 0.3$, let us now pick $\alpha_1 = 0.95$, which locates the coordinates in the full-garbling region (colored in green). In this case, in equilibrium, $\sigma_N^{1*} = 0$. For the new conjectured strategy profile $(\sigma_W^1, \sigma_W^2; \sigma_N^1, \sigma_N^2) = (1, 1; 0, 1)$, proceeding as before, we obtain $\Delta_{r,N}^{1,RP} = -0.004 < 0$, $\Delta_{r,N}^{2,RP} = 0.022 > 0$, $\Delta_{r,W}^{1,AP} = 0.072 > 0$, and $\Delta_{r,W}^{2,AP} = 0.211 > 0$; which corresponds indeed with a fully-garbling equilibrium.

Appendix B. Supplementary data

Supplementary material related to this article can be found online at <https://doi.org/10.1016/j.eurocorev.2024.104940>.

References

- Andina-Díaz, Ascensión, García-Martínez, José A., 2020. Reputation and news suppression in the media industry. *Games Econom. Behav.* 123, 240–271.
- Andina-Díaz, Ascensión, García-Martínez, José A., 2023. Reputation and perverse transparency under two concerns. *Eur. J. Political Econ.* 79, 102439.
- Ashworth, Scott, Shotts, Kenneth, 2011. Challengers, democratic contestation, and electoral accountability. Working paper, SSRN.
- Avery, Christopher N., Chevalier, Judith A., 1999. Herding over the career. *Econom. Lett.* 63, 327–333.
- Banerjee, Sanjay, 2021. Does competition improve analysts' forecast informativeness? *Manage. Sci.* 67 (5), 3219–3238.
- Boudreau, Kevin J., Guinan, Eva C., Lakhaniand, Karim R., Riedl, Christoph, 2016. Looking across and looking beyond the knowledge frontier: Intellectual distance, novelty, and resource allocation in science. *Manage. Sci.* 62 (10), 2765–2783.
- Bourjade, Sylvain, Jullien, Bruno, 2011. The roles of reputation and transparency on the behavior of biased experts. *Rand J. Econ.* 42 (3), 575–594.
- Buisseret, Peter, 2016. "Together or apart"? On joint versus separate electoral accountability. *J. Politics* 78 (2), 542–556.
- Canes-Wrone, Brandice, Herron, Michael C., Shotts, Kenneth W., 2001. Leardhip and pandering: A theory of executive policymaking. *Am. J. Political Sci.* 45 (3), 532–550.

²⁹ These other players are the stronger expert of wise type and the two types of the weaker expert.

- Charness, Gary, Masclet, David, Villeval, Marie Claire, 2013. The dark side of competition for status. *Manage. Sci.* 60 (1), 38–55.
- Charness, Gary, Rabin, Matthew, 2002. Understanding social preferences with simple tests. *Q. J. Econ.* 117 (3), 817–869.
- Corchón, Luis, Serena, Marco, 2018. In: Corchón, Luis, Marini, Marco A. (Eds.), *Handbook of Game Theory and Industrial Organization*, Vol II. Edward Elgar.
- Duesenberry, James S., 1949. *Income, Saving and the Theory of Consumer Behavior*. Harvard University Press, Cambridge, Massachusetts.
- Edelman, Benjamin, Larkin, Ian, 2015. Social comparisons and deception across workplace hierarchies: Field and experimental evidence. *Organ. Sci.* 26 (1), 78–98.
- Effinger, Matthias R., Polborn, Matthias K., 2001. Herding and anti-herding: A model of reputational differentiation. *Eur. Econ. Rev.* 45 (3), 385–403.
- Fehr, Ernst, Charness, Gary, 2024. Social preferences: Fundamental characteristics and economic consequences. Working paper, IZA DP No. 16200.
- Fehr, Ernst, Schmidt, Klaus M., 1999. A theory of fairness, competition, and cooperation. *Q. J. Econ.* 114 (3), 817–868.
- Festinger, Leon, 1954. A theory of social comparison processes. *Hum. Relations* 7, 117–140.
- Fox, Justin, Van Weelden, Richard, 2010. Partisanship and the effectiveness of oversight. *J. Public Econ.* 94, 674–687.
- Fox, Justin, Van Weelden, Richard, 2012. Costly transparency. *J. Public Econ.* 96 (1), 142–150.
- Frank, Robert H., 1985. *Choosing the Right Pond: Human Behavior and the Quest for Status*. Oxford University Press, New York.
- Gentzkow, Matthew, Shapiro, Jesse M., 2006. Media bias and reputation. *J. Polit. Econ.* 114 (2), 280–316.
- Gibbons, Frederick X., Buunk, Bram P., 1999. Individual differences in social comparison: Development of a scale of social comparison orientation. *J. Pers. Soc. Psychol.* 76, 129–142.
- Glazer, Amihair, 2007. *Strategies of the political opposition*. Working paper, University of California-Irvine.
- Green, Jerry R., Stokey, Nancy L., 1983. A comparison of tournaments and contracts. *J. Polit. Econ.* 91 (3), 349–364.
- Harbring, Christine, Irlenbusch, Bernd, Krakel, Matthias, Selten, Reinhard, 2007. Sabotage in corporate contests - An experimental analysis. *Int. J. Econ. Bus.* 14 (3), 367–392.
- Hirsch, Fred, 1976. *Social Limits to Growth*. Harvard University Press, Cambridge, Massachusetts.
- Johnson, Wayne, Proudfoot, Devon, 2024. Greater variability in judgements of the value of novel ideas. *Nat. Hum. Behav.* 8 (3), 1–9.
- Jordan, Christian H., Spencer, Steven J., Zanna, Mark P., Hoshino-Browne, Etsuko, Correll, Joshua, 2003. Secure and defensive high self-esteem. *J. Pers. Soc. Psychol.* 85 (5), 969–978.
- Kreps, David M., Milgrom, Paul, Roberts, John, Wilson, Robert, 1982. Rational cooperation in the finitely repeated prisoners' dilemma. *J. Econ. Theory* 27 (2), 245–252.
- Lazear, Edward P., 1981. Agency, earnings profiles, productivity, and hours restrictions. *Amer. Econ. Rev.* 71 (4), 606–620.
- Levy, Gilat, 2004. Anti-herding and strategic consultation. *Eur. Econ. Rev.* 48 (3), 503–525.
- Levy, Gilat, 2007. Decision making in committees: Transparency, reputation, and voting rules. *Amer. Econ. Rev.* 97 (1), 150–168.
- Li, Ming, Madarász, Kristóf, 2008. When mandatory disclosure hurts: Expert advice and conflicting interests. *J. Econom. Theory* 139 (1), 47–74.
- Lichtendahl, Kenneth C., Grushka-Cockayne, Yael, Pfeifer, Phillip E., 2013. The wisdom of competitive crowds. *Oper. Res.* 61 (6), 1383–1398.
- Nalebuff, Barry J., Stiglitz, Joseph E., 1983. Prizes and incentives: Towards a general theory of compensation and competition. *Bell J. Econ.* 14 (1), 21–43.
- Niiya, Yu, Jiang, Tao, Yakin, Syamil, 2021. Compassionate goals predict greater and clearer dissent expression to ingroups through collectively oriented motives in Japan and the U.S.. *J. Res. Personal.* 90, 104057.
- Ottaviani, Marco, Sørensen, Peter N., 2001. Information aggregation in debate: Who should speak first? *J. Public Econ.* 81 (3), 393–421.
- Ottaviani, Marco, Sørensen, Peter N., 2006a. Professional advice. *J. Econom. Theory* 126 (1), 120–142.
- Ottaviani, Marco, Sørensen, Peter N., 2006b. Reputational cheap talk. *Rand J. Econ.* 37 (1), 155–175.
- Ottaviani, Marco, Sørensen, Peter N., 2006c. The strategy of professional forecasting. *J. Financ. Econ.* 81 (2), 441–466.
- Panova, Elena, 2010. *Confirmatory news*. Working paper.
- Pérez, Antonio-Filiu Franco, 2005. Las visionarias variaciones del conde de aranda respecto del “problema americano” (1781-1786). *Cuadernos de Estudios Del Siglo XVIII* 15, 65–93.
- Prat, Andrea, 2005. The wrong kind of transparency. *Amer. Econ. Rev.* 95 (3), 862–877.
- Reiss, Julian, 2020. Why do experts disagree? *Crit. Rev.* 32 (1–3), 218–241.
- Tullock, Gordon, 1980. In: J.M. Buchanan, R.D. Tollison, Tullock, G. (Eds.), *Towards a Theory of a Rent-Seeking Society*. Texas AM University Press.
- Vandenbroucke, Lucien S., 1984. Anatomy of a failure: The decision to land at the bay of pigs. *Political Sci. Q.* 99 (3), 471–491.
- White, Judith B., Langer, Ellen J., Yariv, Leeat, IV, John C. Welch, 2006. Frequent social comparisons and destructive emotions and behaviors: The dark side of social comparisons. *J. Adult Dev.* 13 (1), 36–44.