

Financial Advisor, Private Incentives, and Conflict Uncertainty

Yixuan Li

A dissertation submitted in partial fulfillment
of the requirements for the degree of
Master of Philosophy
of
University College London.

School of Management and Department of Economics
University College London

September 25, 2025

I, Yixuan Li, confirm that the work presented in this thesis is my own. Where information has been derived from other sources, I confirm that this has been indicated in the work.

Abstract

This thesis examines how a monopolistic financial advisor with conflicts of interest chooses the precision of information sold to a client when that precision is unverifiable. A two-stage game is developed with an informed advisor, a client, a market maker and a noise trader in a modified Kyle (1985) framework, trading a single binary-valued asset. The advisor earns revenue from both the sale of a signal to the client and from a separate service line whose payoff relates to the asset and may move with or against the client's trade.

Two scenarios are analysed. In Scenario 1 the advisor observes the magnitude and direction of the conflict when pricing the signal. Either the equilibrium fee and signal precision are made contingent on conflict size, or a single fee and signal distribution is offered that does not vary with conflict. Revealing the conflict does not guarantee full information, as it discloses only the relative, not absolute, scale of the private benefit. In Scenario 2 the advisor prices using only a distribution of the conflict, but observes the realisation when supplying the signal. Because the realised conflict can help or hurt the advisor's outside payoff, uncertainty dilutes the incentive to manipulate and changes the trade-off between precision revenue and the outside payoff. Multiple equilibria arise, including pooled-fee outcomes in which realised precision may be the same across conflict states or state-dependent. Over relevant parameter ranges this yields a non-monotonic relationship between conflict and realised precision, so larger conflict does not always degrade advice quality.

The thesis extends discussion of biased advisors in informed trading models. Policy implications for disclosure and monitoring conflicts are also explored.

Impact Statement

This research analyses how conflicts of interest affect the quality of information that financial advisors provide when the precision of that information cannot be verified. The model isolates two situations faced in practice: advisors who know their conflict at the point of giving advice, and advisors who know only a distribution of the conflict that is realised later.

The findings speak to regulators of advisor conduct. Disclosure alone does not guarantee full information quality: when precision is unverifiable, a declared conflict can still leave room for strategic shading. The usual expectation is that larger realised conflicts mean lower precision. The model offers a counterpoint: when the conflict is uncertain at the advice stage, reputational and repeat-business incentives can sustain, and sometimes even increase, precision, even if the realised conflict later turns out somewhat larger. In practical terms, advisors may switch between advice strategies as the conflict is resolved, so assuming any conflict must always reduce quality can be misleading. Disparities between possible conflict size that the advisor is facing also affects information quality, which might reflect in real world cases.

Advisors commonly operate alongside market-making, underwriting, asset-management or research functions, so some conflict is unavoidable and robust information barriers between service lines can be difficult to maintain. The model shows that limited and well-monitored conflicts need not lower information quality; under uncertainty, reputational incentives can sustain information precision even if realised conflicts become slightly larger. From a policy perspective, it is helpful to monitor and limit realised conflicts, use market incentives to create reputational

incentives for accurate advice (especially when conflict is uncertain), determine appropriate cutoff for disclosure or auditable records that raise the cost of shading. These insights also inform how to design and implement mandatory disclosure and commission rules for financial advisors, including when to combine tools. Identifying whether the conflict is known or uncertain when advice is quoted matters, because the effects of disclosure and limits on conflict size can differ across advice markets.

The work also highlights limits: it is a stylised theoretical model with binary states and assumes rational pricing by a market maker. Further empirical work that varies conflict uncertainty could test the mechanisms identified here.

Acknowledgements

The substantive work was completed before May 2024 under earlier supervision; subsequent steps have been administrative. I acknowledge and am grateful for the support provided by Dr Colin Fisher and Dr Saleem Bahaj in coordinating the examination. I would also like to thank all those, professional and personal, whose advice and effort helped me to secure the opportunity to submit this thesis. My gratitude extends as well to my peers for their feedback in our seminar series and ongoing discussions. My heartfelt thanks go to my family and friends for their patience and support throughout this transition.

“All human wisdom is contained in these two words, ‘Wait and hope.’”

— *Alexandre Dumas, The Count of Monte Cristo*

Contents

1	Introduction	10
2	Model	20
2.1	Setup	20
2.1.1	Players, states and timeline	20
2.1.2	Information game actions	20
2.1.3	Asset trading game actions	21
2.1.4	Payoffs	22
2.1.5	Equilibrium Definition	23
3	Equilibrium	24
3.1	Strategies, solution method, and equilibrium types	24
3.2	A simple benchmark: no private incentives	31
4	Solving the Baseline Model	33
4.1	Baseline, case 1: semi-pool fee and semi-pool precision	33
4.2	Baseline, case 2: fully pool fee and fully pool precision	36
4.3	Baseline, case 3/4: fully pool fee and semi-pool/cross-pool precision	37
5	Advisor Type Uncertainty: A Modified Problem	41
5.1	Type uncertainty, case 2: fully pool fee and fully pool precision	42
5.2	Type Uncertainty, Case 3/4: Fully Pool Fee and Semi-Pool/Cross-Pool Precision	43

5.2.1	Case 3.1: Truthtelling High Type and Partial Truthtelling Low Type	45
5.2.2	Case 3.2: Truthtelling High Type and Babbling Low Type	45
5.3	Understanding the Equilibria	46
5.3.1	Summary of Analytical Results	47
5.3.2	Substitutability of Advisor's Income Sources	48
5.3.3	Equilibria Comparisons	51
5.3.4	Resolve the Conflict Uncertainty or Not?	53
6	Conclusion	56
	Appendices	57
A	Solve the price-setting process	58
B	Proof for Lemma 1	64
C	Appendices for analysis sections	65
	Bibliography	66

List of Figures

2.1	Model timeline	20
3.1	Graphical definitions of <i>Semi-pool</i> , <i>Fully pool</i> , and <i>Cross-pool</i> for $R_{\theta,k}$. Identical colours denote equality. For precision, replace R with α	28
4.1	Conflict-revealing Equilibria: Expected payoff as a function of conflict k	34
5.1	Equilibria with type uncertainties: Interior solutions	46
5.2	An example with $k_H = -\frac{1}{2}$. Dashed line: α_{k_H} . Brown line: α_{k_L}	50
C.1	Region When Fully Pool Equilibrium Achieves a Higher Advisor Utility	67

Chapter 1

Introduction

The financial advice market involves an advisor providing financial planning advice and trade execution services for clients. It is widely acknowledged that this principal-agent relationship is not free of conflict-of-interest, and the advisors may benefit from actions hidden from the client. This paper examines two aspects of conflicts in this moral hazard problem.

First, consider a scenario where advisors have private benefits that are directly or indirectly tied to themselves or their employer's profit. To maximise their own gains, an advisor may adjust the level or content of information according to personal interests, especially when in the absence of regulation. It is particularly worrisome when these alternative sources of benefits are not easily observable or verifiable, and regulatory costs on monitoring and compliance are high.

This material conflict-of-interest problem is relevant to many real-world scenarios. Empirical literature establishes that there are agency problems related to conflict-of-interest through commissions, bank profits and misaligned product recommendations. This situation creates a negative incentive for the agent in terms of information provision, resulting in a loss of investor welfare (e.g. Hackethal et al. (2012); Hoechle et al. (2018); Chalmers and Reuter (2020); Egan (2019)). Regulatory bodies have also recently drawn attention to the principal-agent cross-trade problem with investment advisors to benefit from matching clients' trades at the expense of the client rather than aiming for best execution.¹ Concerns also exist

¹See the reports: <https://www.sec.gov/files/OCIE%20Risk%20Alert%20-%20Principal%20and%20Agency%20Cross%20Trading.pdf> and <https://www.>

about the coexisting activities of investment banks in providing financial advice and engaging in market trading since 2008. This includes service lines such as discretionary trading arms, market-making and trading book adjustment, and their potential to manipulate markets. Since the financial crisis, the United States introduced regulations such as the Volcker rule that essentially banned proprietary trading by banks and their linkages to hedge funds. Nevertheless, regulatory bodies encountered difficulties in assessing the effects and effectiveness of the current regulations (Duffie (2012); Bao et al. (2018)). In 2020, the Federal Reserve eased certain measures and granted exemptions for specific market activities by banks with advisory, trading delegation or market-making capacities, while regulatory actions monitored the prohibited activities of irregular size and issued penalties. One cited reason is the fear of hindering banks' other services mentioned, and these functionalities could be important for financial market stability. One aim of this paper is to provide insights into the scope of the agency problems mentioned and assess the effectiveness of regulation from a planner's perspective.

Second, another related and heavily debated issue in the financial market centres on the type and transparency of fees that an agent should collect for issuing recommendations or for asset management. Many fee structures prevail in the market; widespread concerns exist regarding whether some implicit or conflicting incentives in fees affect information efficiency. Examples include assets-under-management fees (AUM) charged by investment advisors or when advisory fees also include client transaction costs and are charged as a package rather than as separate fees (regulation on the latter has been done in Europe; for details, see MiFID II). Whether these fees directly reveal conflict of interest of advisors is another question. From the fee-unbundling regulation in Europe, there are indeed concerns that a client might pay the advisory fee without fully realising the underlying incentives of their advisors. In general, using explicit regulatory tools to ensure mandatory disclosure of conflict when advising the client can be helpful. However, whether such disclosure rules of conflicts (Li and Madarász (2008); Stoughton et al. (2011); De

[sec.gov/files/fixed-income-principal-and-cross-trades-risk-alert.pdf](https://www.sec.gov/files/fixed-income-principal-and-cross-trades-risk-alert.pdf)

Moragas (2022)), whether through form of fees like kickbacks and commissions or general forms of conflict, are indeed beneficial for information provision is another question that remains unconcluded (Li and Tiwari (2009); Ma et al. (2019)). Also, an advisor's conflict can occur at time either before or after an advisory relationship has been formed; in the latter case, mandatory disclosure or fee unbundling might not be the most effective method because of uncertainty about the conflict. (Li and Madarász (2008)) This model seeks to provide an alternative explanation to reconcile the wide range of fees that exist, assess the effect of regulation, and understand unclear empirical predictions of its impact on advisors' investment performance.

Importantly, in both of the conflicts mentioned above, the advisor's information quality cannot be directly verified when the client trades in the market; otherwise, the agency problem would be largely resolved. Typically, one expects (and it is commonly assumed in theoretical modelling) to observe information quality at the same time as cost quotes. However, many factors can render quality an unobservable variable to the public and hence, give the advisor manipulation potential. In this context, I feature transparency issues and misaligned private incentives of the advisor as two main channels that subsequently affect information quality. When the client observes the committed information quality, the advisor does not have full freedom in setting the signal structure contingent on their private benefit, and sometimes this would force the advisor to disclose their private benefit even in the absence of regulation. By contrast, the advisor's private-incentive realisation can play a confounding signalling role in the equilibrium signal structures, and the information fee does not fully signal the payoff-irrelevant state. Consequently, the client receives signals that are also contingent on the payoff-irrelevant state and resulting in the advisor hiding what they know about the financial asset.

In the project, I studied in a stylised way the extent of changes in signal probabilities that an advisor sent to the client regarding a binary-valued financial asset when the advisor has an additional private benefit (conflict), with binary, stochastic scale parameter, and linked to ex post efficiency. I first consider the scale parameter of such private benefit is known to the advisor when deciding the signal and

fee structures, where each of the variables can be contingent on both asset value and conflict scale, and signals and fees each take no more than four cases. In this baseline, I characterise the form of signal structure the advisor would adopt. With unobserved signal precision and private state, the advisor can potentially make the signals they send to the clients contingent on both states, as long as the publicly known information (e.g., signal price) does not enable information free-riding on the asset value. Then, signal structures and the resultant fees can take many forms.

Apart from two easy-to-spot equilibria featuring full pooling with respect to both signal and fee structures, or revealing private states through signal price, two other signal structures exist in this model. These structures send asymmetric signals with respect to the private states or jointly with both states, but fees fully pool. The common feature is that advisors choose not to signal their private state. However, under the baseline, only the former two types of equilibria persist. The reason is as follows: When an advisor knows the relative strength of the conflict, and cares about the realised signal performance, her private benefit is only contingent on the signal precision conditional on such conflict state parameter. However, the fee charged in a no-signalling scenario is jointly determined by realised and unrealised signal performance, and the client does not fully update their beliefs on the realised conflict values. This creates a mismatch in incentives to buy the signal, as the advisor finds it optimal to charge the maximum fee conditional on the unrealised signals, but the client only wants to buy when expected trading profit, corresponding to the expected value of information received, is weakly higher than the fee. In this case, the client knows that the fee includes a component from unrealised precision, but the realised precision is lower when the advisor charges the fee comparing to maximum trading profit. Then, the advisor ends up with either a conflict-revealing equilibrium, which signals her private incentive in the fee, when there is sufficient distance between the two types or both types prefer to sell partial truth due to conflict; or, a conflict-hiding equilibrium, which does not signal the incentive, but provides full information content or drops out of the market. Generally, if the conflict scale is not sufficiently large, an advisor is able to tell what is known to the client, and full

incentive alignment, i.e. private benefit increases with information precision, is not a necessary condition. However, revealing a conflict of interest does not necessarily mean that the client will receive full information content.

I then analyse how uncertainty about advisors' conflicts affects equilibrium outcomes. Due to uncertainty, advisors have all conflict-hiding equilibria at their disposal. The crucial difference in this modified model is that her private value changes to an ex ante measure of realised signal performance, weighted by the values of the conflict scale parameter. The previous incentive mismatch disappeared due to this private value in ex ante form, and thus equilibria previously eliminated are restored here. There are at most three different cases an advisor can face: providing different realised precision to the client once conflict scale (i.e. advisor's type) is realised, with the high type telling the client what she knows. The remaining case is where the advisor pools fully, and both types can tell partial truth. When advisor shifts between one equilibrium and another in this multiple equilibria scenario, curious results appear. The low type advisor has a non-monotonic and discontinuous decision for information precision with respect to the conflict scale parameter, contrary to the conventional wisdom that restricting conflicts for advisors helps information provision. The reason traces to the sensitivity of equilibrium solutions to the advisor's conflict scale, which determines the domains for equilibrium existence and the fully pooled equilibrium is the least sensitive. Then, the advisor moves from an equilibrium where the high type tells full information, but the low type tells nothing, to another equilibrium where precisions are fully pooled but types provide some information. It is also possible that within some parameter combinations, multiple equilibria exist. As a final note, contrary to the baseline, there is no clear conclusion that the advisor necessarily recoups higher utility at the expense of client's signal performance. The question of resolving uncertainty in conflicts would also depend on the magnitude and dispersion of the advisor's conflict scales.

I then discuss regulatory tools to promote higher information quality under conflict uncertainty. Based on the model, an alternative explanation exists for the fact that, although an advisor's conflict might be large, and/or compensation struc-

tures argued to hide conflict are often preferred by the advisor if the industry faces less regulation, the effects on signal probabilities are not necessarily negative. Providing more aligned incentives to advisors is important, followed by examples to improve reputation costs for advisors, fines for advisors with large conflict, or industry awards to advisors who provides higher-quality information, resulting in client success. However, the impact of such actions would depend greatly on whether conflicts of interest are realised for advisors prior to forming an advisory relationship with the client. It is possible to trigger more communication by simply providing such incentives to the relatively good type of advisors, when such type is unrealised to advisors prior to charging the fee. It is also possible to induce limited, or negative benefits from such policy if restrictions are only linked only to punishing when an advisor turns out to have relatively larger conflicts, due to multiple equilibria. To reduce this problem, fully aligning incentives with the client helps. The results from the model call for a cautionary approach to tighter intervention within advisory markets, while echoing empirical and theoretical evidence on the relationship between an agent's compensation and performance.

The structure of the paper is as follows. The rest of this section reviews the related literature. Sections 2 and 3 set out the model and equilibrium strategies. Section 4 discusses the baseline model and its implications. Section 5 extends the baseline model and presents a case with uncertainty about advisor's bias. Section 6 concludes. The appendix includes relevant formal proofs and supplementary material.

Related literature. This paper joins the discussion of information transmission with conflicts of interest, with a particular element of selling information. Early papers on standard information sales problems with a monopolistic advisor consider the agent's risk-sharing benefits of the information-selling revenue, for which the precision of information is observable. (Admati and Pfleiderer (1986, 1988), Veldkamp (2006), Cespa (2008), García and Sangiorgi (2011), to name a few) Although information sales provide risk-sharing incentives, multiple factors prevent full information sharing between the advisor and the client. Optimally adding photocopying

or personalised noise to avoid information free-riding is a natural outcome. In a risk-averse setting, it is also possible to reduce information production.

I modify some assumptions in the above examples to show other channels of lower signal precision from conflict of interest. One is the aligned objective function between the advisor and the client, i.e. risk-sharing incentives dictate the advisor's behaviour, and the advisor knows the true asset state, dimension one, and there are no other payoff-irrelevant states. Another is that signal precision is observable to the client so that the expected interim trading gains only reflect signal precision as the ex ante value of information. A separate thread of literature on informed trader precision uncertainty and unobservable information acquisition exists (e.g. Chakraborty and Yılmaz (2004), Banerjee and Green (2015), Banerjee and Breon-Drish (2020), Xiong and Yang (2023)). However, in my model, the focus is mainly on its impacts to hide the conflict of interest by the advisor, resulting in manipulation, in terms of hiding the asset-value information known to advisor, and contingent on a payoff-irrelevant state.

The paper broadly fits into the topic of communication with conflicts of interest. Early works are based on a cheap-talk framework and primarily focus on disclosure (e.g. Benabou and Laroque (1992), Morgan and Stocken (2003), Li and Madarász (2008), Gesche (2021)). In particular, the last example develops the question in which the conflict of interest becomes uncertain to the principal. I made two distinctions in modelling. Different from these models, I develop the problem in a trading scenario to explicitly model the value of information sold, represented by the fee, as in information sales literature, rather than disclosed information value. This means I solve the equilibrium by demonstrating the signalling structure of the advisor's private benefit state, but not an information-partitioned equilibrium with respect to the asset state that the client cares about in the cheap-talk literature. The advisor also does not have a preference with respect to the asset-value realisations, explicitly in payoff function. Both approaches, however, provide insights into when the game discourages truthtelling, given that the advisor is biased. However, in my modified model, it is possible to show that information quality is not necessarily

monotonic to the conflict when advisor's bias becomes uncertain, established first in Li and Madarász (2008). They also shed light on why mandatory disclosure of an advisory firm might not be optimal for information quality.

More recent papers have evolved to document various and distinct sources of conflict of interest in information production and methods to improve information efficiency. Broadly, they address situations where advisors have reputational concerns or instead, enhancing information quality may not align with the advisor's interests. The works of Inderst and Ottaviani look at the problem from an optimal contract perspective for third-party agents tasked with both product recommendation and selling roles. They identified the agency costs to solve this incentive problem of the separate information mediator (Inderst and Ottaviani (2009)). They went on to investigate the effect of a product producer's hidden commissions on information intermediaries (Inderst and Ottaviani (2012a)) and emphasise the importance of financial literacy. Inderst and Ottaviani also reviewed the policies in place (Inderst and Ottaviani (2012b)). Some other theoretical and empirical evidence on such conflict problem include, e.g. Stoughton et al. (2011), Bhattacharya et al. (2012); Hackethal et al. (2012); Calcagno and Monticone (2015); Chalmers and Reuter (2020); however, these papers also highlight caveats when promoting policies to restrict conflicts, as it can be ineffective. These are joined with empirical and experimental evidence to wider range of topics with disclosure provide mixed effect on information quality available in the market. (Kartal and Tremewan (2018); Ismayilov and Potters (2013); Behnk et al. (2014)) For example, in Chang and Szydlowski (2020), as a related work to Inderst and Ottaviani (2009), the advisor can only recoup profit upon client execution, but the advisors in their context have multiple types and different to this paper, multiple advisors exist and compete for customers. Another difference is that the type relates to information on the asset rather than to an advisor's private benefit. With multiplicity in the advisor or clients' types, regulation to improve the lower bound of the advisor's signal precision helps the advisor exhibit complementarity in information production, and one of the paths could be improving the client's financial literacy. However, without taking into ac-

count equilibrium responses of type multiplicity, rules that directly target fee level without considering signal precision level in equilibrium do not improve welfare. In a very different context, Malenko and Malenko (2019) analyses corporate voting and proxy advisors, where competition for a private information source from the principal can prevent certain cases of information improvement by crowding out other information acquisition channels. Similarly, there is a threshold for the advisor to be beneficial for the voter's information quality. Reputation costs are one of the common-sense solutions to information efficiency. However, papers establish that reputation itself is not sufficient for the advisor to self-regulate (Ottaviani and Sørensen (2006)). As other countermeasures to the moral hazard cost, Bolton et al. (2007) and Bolton et al. (2012) discuss the issue of reputation costs and competition among advisors in various contexts, respectively, in agents misselling financial products (like in Inderst and Ottaviani (2009)) in direct price competition among information providers, and with credit ratings, and demonstrate mixed impacts of these two factors in improving information efficiency.

The way to model private benefit also connects to the feedback effects literature. (Dow and Gorton (1997); Goldstein and Guembel (2008); Goldstein et al. (2013); Angeletos et al. (2022); Dow et al. (2017)) This literature focuses on a feedback loop between a firm's investment decision which maps into ex post cash flows, and firm's value in secondary market trading before such investment, where a firm can learn from valuation influenced by potential insider trading, and the impact on investment decisions and ex post payoffs. There are similarities when an advisory firm's private value relates to ex post outcome of signal performance, or in general, ex post outcomes relates to asset market that trading happens after information is acquired. The advisor's strategies, i.e. information precision contingent on private value scales and choice of fees, impacts on advisor's utility and private value but also reflects the advisor's private value scale. In a way, prices play roles very similar to those feedback effect papers, and so mechanisms like strategic complementarity, including those in information production scenarios and multiple equilibria, arises. A very good example combining these is Dow et al. (2017), focus on both feedback

effects and information production, and also features a multiple equilibria mechanism with respect to information quality. However, learning in this model happens not only in the asset value through market prices, but also in the advisor's private value, which occurs in the information market, making this a joint, and more complex, learning problem of an informed trader and a market maker in a microstructure model. In my simplified problem, however, as price does not reflect or impact the ex post realisation of signal performance, the argument in Goldstein and Guembel (2008) that an uninformed speculator can make profitable trades cannot hold. A conjecture is that extending the model with a slightly complicated noise trader might realise the mixed strategy pattern in the mentioned paper.

As a final remark, the work loosely links to works that discuss the optimal form of compensation for advisors and its link to the principal's benefit, with a focus on the asset management industry. The literature addresses widespread concerns about the negative relationships between compensation structures and managerial incentives or investment performance. Theoretical insights are offered by papers such as Starks (1987), Admati and Pfleiderer (1997), Ross (2004), Li and Tiwari (2009), Cuoco and Kaniel (2011), and these theoretical perspectives find empirical support as well. Indeed, as the extent of the agency problem increases, agents tend to prefer less flat compensation structures. However, the relationship between compensation structure and investment performance remains somewhat unclear; optimal contracting at equilibrium means the expected level of compensation matters rather than its form, and non-linear contract can be optimal. (Li and Tiwari (2009), Ma et al. (2019)) However, there is also evidence of how referencing other benchmarks might affect investment performance (Admati and Pfleiderer (1997), Sotes-Paladino and Zapatero (2022), Li and Wu (2019)). Changes to an advisor's gambling attitude also exhibit non-monotonic patterns (Golec and Starks (2004)). This paper investigates an alternative uninformed response channel as motivation for the fee structure. It shows the impact of this uninformed decision rule on the agency problem and the interaction with the presence of the advisor's private incentive as an additional term in the direct utility.

Chapter 2

Model

2.1 Setup

2.1.1 Players, states and timeline

There are four players: a financial advisor, a client, a market maker and an uninformed traders. There are four dates, $t = \{0, 1, 2, 3\}$. The timeline of the game is as follows. At time 0, nature moves and draws two independent random variables: the asset value $\theta \in \{0, 1\}$ with prior $\mathbb{P}(\theta = 0) = \frac{1}{2}$ and $k \in \{k_L, k_H\}$ being the additional incentives, with prior $\mathbb{P}(k_L) = \frac{1}{2}$. At time 1, the client and the advisor interact in an information game. At time 2, the asset is traded in a competitive market. After trading at time 2, players' payoffs are realised.

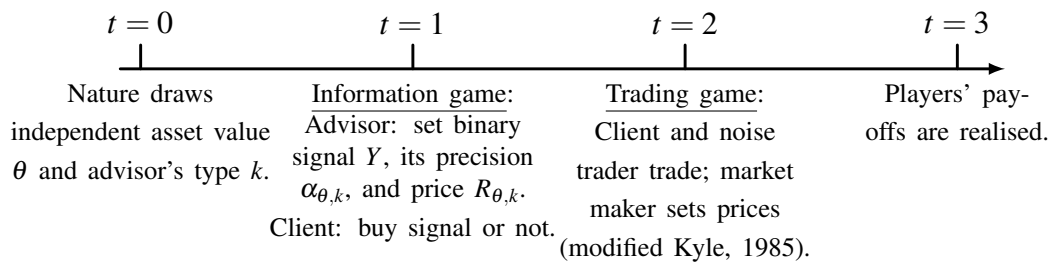


Figure 2.1: Model timeline

2.1.2 Information game actions

At time 1, advisor designs the advice to the client in the form of a private signal realisation (signal) $Y \in \{0, 1\}$ revealed at time 2, and its underlying probability

(precision)

$$\alpha_{\theta,k} := \mathbb{P}(Y = \theta | \theta, k)$$

and sets a fee $R_{\theta,k}$, which would be publically known to all market participants. Her information F_A at time 1 is $F_A = \{\theta, k\}$, i.e. the advisor knows the private incentive realisation when enters the information market, and have the option for signalling. Corresponding examples would be existing deals (e.g. M&A) of the firm which serves as potential conflict of interests. Advisor commits to the actions fully. Assume that advisor presents the price $R_{\theta,k}$ as a public and take-it-or-leave-it offer. This full bargaining power assumption is for analysis simplicity and relax it would not hurt the results. The client has no other source of information about the asset. The client observes signal price $R_{\theta,k}$ and determines whether to buy the signal or not.

2.1.3 Asset trading game actions

The price-setting process follow a modified Kyle (1985) setup with fixed trading size and binary asset value. There is one exogenous noise trader. Assume that the noise trader has no information on the asset state, trades for non-information reasons and randomly buy or sell one unit, with exogenous buy probability $\mathbb{P}(U = 1) = \varepsilon$. Then, the uninformed trader's actions are $U \in \{-1, 1\}$. If bought the signal in period 1, the client observes Y before asset market opens, and stands as an informed trader. If not bought the signals, the client knows nothing about asset state θ . The informed trader's actions are trading unit $I \in \{-1, 0, 1\}$, probabilities $\{\mu_{Y,R}\}$ with which client follows signal Y :

$$\mu_{1,R} = \mathbb{P}(I = 1 | Y = 1, R_{\theta,k}), \quad \mu_{0,R} = \mathbb{P}(I = -1 | Y = 0, R_{\theta,k})$$

and probabilities with which the client chooses to no trade with information

$$\mathbb{P}(I = 0 | Y, R_{\theta,k}).$$

Traders submit market orders. A market maker (MM) sets prices and provides liquidity. MM observes the aggregate order flow $d = I + U$ and provides price schedule, p_d , conditional on information she knows, $F_{MM} = \{d, R_{\theta,k}\}$.

2.1.4 Payoffs

Assume that agents are risk-neutral. The advisor recoups payoff EU_A from both subgames. She receives $R_{\theta,k}$ if client buys the signal, and receives an additional private benefit from an unmodelled source in time 2 as $k\pi Q$. Define

$$\pi := \mathbb{E}(I(\theta - p)|Y, R_{\theta,k})$$

as the client's expected trading profit, i.e. the value of information, when client does buy the signal. The function Q represents the advisor caring about the signal informativeness they bring to the ex-post asset market outcome. To measure the effect of their signals on these ex-post outcomes using a verifiable function, and represent potential conflicts between advisor and client, I use nonnegative realised trading profit and $Q = \mathbb{1}(I(\theta - p) \geq 0)$ which is equivalent to $Q = \mathbb{1}(Y = \theta)$. k represents the uncertain scale and direction of such private benefit of the firm relative to the expected trading profit π . Then, sign of k indicates whether the advisor's incentive is aligned with the client or not and this private benefit represents a direct source for potential conflict of interest (CoI). If $k > 0$, advisor has solely performance concerns, and hence fully aligned incentive. In this case, $k\pi$ is benefit of telling client the truth. Similarly, if $k < 0$, advisor has misaligned incentive and lying gives private benefit gains; $|k|\pi$ represents advisor's benefit of lying. A proxy for such private benefit would be cross-trading revenue or trust-building benefits. Advisor's expected payoff is shown below.

$$EU_A = \underbrace{R_{\theta,k} \times \mathbb{1}(\text{client buy signal})}_{\text{time 1}} + \underbrace{\mathbb{E}[k\pi Q|F_A]}_{\substack{\text{unmodelled other service lines,} \\ \text{potential conflict of interest (CoI)}}$$

In reduced form, the expected private benefit of client is

$$\mathbb{E}[kQ\pi|F_A] = k\pi\alpha_{\theta,k}.$$

The client receives the net trading profit

$$S = \pi - R_{\theta,k}.$$

The market maker gets zero expected profit from the price-setting process under Kyle-style setup.

2.1.5 Equilibrium Definition

The solution concept is a Perfect Bayesian Equilibrium (PBE). In equilibrium:

1. **Advisor:** chooses signal precision $\alpha_{\theta,k} := \mathbb{P}(Y = \theta \mid \theta, k)$ and price $R_{\theta,k}$ at time 1 to maximise expected utility EU_A .
2. **Client:** decides whether to buy the signal at $t = 1$, and trading unit I , signal obeying probabilities $\mu_{Y,R}$ and probabilities $\mathbb{P}(I = 0|Y, R_{\theta,k})$ at $t = 2$ to maximise expected surplus.
3. **MM:** determines a price schedule $p_d \in \{p_{-2}, p_{-1}, p_0, p_1, p_2\}$, conditional on the aggregate order flow d such that the market maker makes zero expected profit.
4. **Client and MM:** form rational expectations on $\alpha_{\theta,k}, k, Y$ as $\hat{\alpha}_{\theta,k}, \hat{k}, \hat{Y}$; **advisor and client** additionally on p_d as $\hat{p}_d$. Beliefs are Bayesian updated.

The game is solved by backward induction.

Chapter 3

Equilibrium

3.1 Strategies, solution method, and equilibrium types

The client chooses the trading strategies $\{I, \mu_{Y,R}, \mathbb{P}(I = 0|Y, R_{\theta,k})\}$ to maximise her objective $\pi - R_{\theta,k}$, by obeying the signal in probability $\mu_{Y,R}$. To consider cases when the client gets zero, consider a breakeven assumption that client in $t = 2$ follows signal, i.e. $\mu_{Y,R} = 1$, whenever indifferent between choices of $\mu_{Y,R} \in [0, 1]$. To see the equilibrium $\mu_{Y,R}$, the client's trading profit needs to be computed given player's strategies. For the market maker, the zero expected profit condition implies that price is the expected value of the asset given what she knows, $p_d = \mathbb{E}(\theta|d, R_{\theta,k}) = \mathbb{P}(\theta = 1|d, R_{\theta,k})$. The advisor is the one behind the informed trader, the client, and supplies the signal Y by pinning down precision $\alpha_{\theta,k}$ at price $R_{\theta,k}$ to maximise objective EU_A .

Before analysing time-2 strategies, I make an assumption to rule out the client's choice of randomisation, and considerably simplify the model. I look at equilibrium that if client did not buy signal in time 1, client would not trade, i.e. $I = 0$ with probability 1. Without information, the client is a strategic uninformed trader in the market. She can either trade by randomising between trading actions, which should generate zero expected profit (Jarrow (1992)), or chooses the assumed action, which gives the same outcome. Then, market maker observes client buy or sell only if she is informed. From this assumption, one would be able to know that if client trades

in the market upon observing the signal realisation Y , the client mimics the uninformed trader and trades one unit. To see that this is indeed the case, one can follow the below logic. The client will not choose the no-trade action upon observing Y as no-trade would provide zero payoff, but by hiding behind the uninformed, the client might earn nonnegative expected payoffs. A formal proof and the exact price schedule are in Appendix A.

At time 2, I compute the prices in the asset market given traders' and the advisor's strategies. By Bayesian updating, taking beliefs of θ and k given $R_{\theta,k}$, which is a realised variable at beginning of time 2¹:

$$\mathbb{P}(\theta = 1 | d, R_{\theta,k}) = \frac{\mathbb{P}(\theta = 1 | R_{\theta,k}) \mathbb{P}(d | \theta = 1, R_{\theta,k})}{\mathbb{P}(\theta = 1 | R_{\theta,k}) \mathbb{P}(d | \theta = 1, R_{\theta,k}) + \mathbb{P}(\theta = 0 | R_{\theta,k}) \mathbb{P}(d | \theta = 0, R_{\theta,k})}$$

When $d \in \{-2, 2\}$, as said, the client, as an informed trader, has strategies revealed to the market maker. At $d = 2$, market maker knows an informed trader is buying and at $d = -2$, market maker knows an informed trader is selling. Then, market maker fully adjusts belief to the informed's belief, and compute prices. Take $d = 2$ as an example:

$$p_2 = \mathbb{P}(\theta = 1 | I = 1, R_{\theta,k})$$

If market maker observes $d \in \{-1, 1\}$, MM knows no informed trader presents and thus prices asset at the prior. At middle price p_0 , market maker is confounded with the informed trader action, even if she knows an informed trader exists in this aggregate order flow. For profit calculation, at $d = \{-2, 2\}$, market maker's belief gives zero expected profit for the informed. Only at $d = 0$, an informed trader can earn positive expected trading profit. In general, given rational expectations on

¹If market maker does not know $R_{\theta,k}$, it will result in a simpler version of the Bayesian updating process and a subset of candidate equilibria that we are facing.

prices $p_d \in \{p_{-2}, p_0, p_2\}$:

$$\begin{aligned}\pi &= \mathbb{E}(I(\theta - p_d) | R_{\theta,k}) \\ &= \sum_{\theta \in \{0,1\}} \mathbb{P}(\theta | R_{\theta,k}) \mathbb{E}(I(\theta - p_d) | \theta, R_{\theta,k})\end{aligned}$$

This expected trading profit is a function of client's trading strategies

$$\{I, \mu_{Y,R}, \mathbb{P}(I = 0 | Y, R_{\theta,k})\}$$

and advisor's strategies

$$\{\alpha_{\theta,k}\}$$

The client's equilibrium $\mu_{Y,R}$ maximises π at $\mu_{Y,R} = 1$, and $I \in \{-1, 1\}$ with probability 1. The intuitive reason is the signal Y is the only information source for the client, and even if the client knows it is a biased signal, follow signal is better than not following, which is a weakly uninformative action. As in Appendix A, this implies that when client maximising trading profit, she wants to maximise the correct trading probabilities, i.e. probabilities of buying when asset value is high and selling when asset value is low. Then, the client follows the signal with probability 1 to maximise those correct trading probabilities.

Then, I move to time 1. Again, set a breakeven assumption to consider non-trivial actions from the advisor: in $t = 1$, client buys signal whenever she is indifferent between buying and not buying. Combined with full bargaining power of the advisor, this allows advisor to purpose the maximum amount of information value π as price $R_{\theta,k}$ and client still accepts this offer. Therefore, client in this case always buy signal.² Recall that advisor set precision $\alpha_{\theta,k}$ at price $R_{\theta,k}$ to maximise objective $EU_A = R_{\theta,k} + k\alpha_{\theta,k}\pi$. The next step is to consider space of $(R_{\theta,k}, \alpha_{\theta,k})$ and back out the possible combinations.

I start with a simple intuition. Without k in the advisor's type space, the advisor

²In case that advisor does not have full bargaining power, there is an upper bound for $R_{\theta,k}$ that advisor can charge. The client's indifference assumption suggests client buys the signal at this upper bound and the remaining argument still holds.

would never contingent the offer on the asset state θ to avoid the risk of information free-riding; that is, the client receives information about θ from observing the fee $R_{\theta,k}$, without actually paying it. However, with a camouflage of a private type k , the situation becomes much more complex. We have four possible outcomes for $(R_{\theta,k}, \alpha_{\theta,k})$ respectively. The free riding reasoning allows ruling out $R_{\theta,k}$ that resulting in posterior belief $\mathbb{P}(\theta = 1 \mid R_{\theta,k}) \in \{0, 1\}$, i.e. asset state value revelation. These include cases like $R_{\theta,k}$ are all different,

$$R_{0,k} \neq R_{0,k'} \neq R_{1,k} \neq R_{1,k'}$$

only one $R_{\theta,k}$ is fully revealing,

$$R_{\theta,k} \neq R_{\theta,k'} \neq R_{\theta',k} = R_{\theta',k'} \quad \forall \theta \neq \theta', k \neq k'$$

and $R_{\theta,k}$ is fully revealing in θ .

$$R_{0,k} = R_{0,k'} \neq R_{1,k} = R_{1,k'}$$

The advisor has remaining options to contingent $R_{\theta,k}$ on k , pool the fee, or jointly (θ, k) . To simplify language, name the three fee cases $R_{\theta,k}$ as semi-pool

$$R_{0,k_L} = R_{1,k_L} \neq R_{0,k_H} = R_{1,k_H}$$

$$R_{\theta,k} = R \quad \forall (\theta, k)$$

and cross-pool respectively.

$$R_{0,k_L} = R_{1,k_H} \neq R_{0,k_H} = R_{1,k_L}$$

Following the reasoning above to outline possible cases in the signal space, semi-pool, fully pool and cross-pool signal precision exists. Consider first symmetric signals with respect to θ , i.e. semi-pool and fully pool case, supplies two possible equilibria: whether to separate the signals with respect to k or not. The

advisor can issue a semi-pool signal, separates signals with respect to k but pooling with respect to θ .

$$\alpha_{\theta_H, k_H} = \alpha_{\theta_L, k_H} \neq \alpha_{\theta_H, k_L} = \alpha_{\theta_L, k_L}$$

Also, a fully pool signal consists of:

$$\alpha_{\theta_H, k_H} = \alpha_{\theta_H, k_L} = \alpha_{\theta_L, k_L} = \alpha_{\theta_L, k_H} = \alpha$$

With the help of the additional state variable k , the advisor is able to send cross-pool signals as asymmetric signals with respect to θ .

$$\alpha_{\theta_H, k_H} = \alpha_{\theta_L, k_L} \neq \alpha_{\theta_H, k_L} = \alpha_{\theta_L, k_H}$$

Further revealing θ will encourage information free-riding; therefore, any other combinations that reveal a pair of θ, k are not considered. For clarity, the patterns for fees $R_{\theta, k}$ are shown below; the same patterns apply to precision $\alpha_{\theta, k}$ by replacing R with α . Identical colours indicate equal fee/precision levels.

Semi-pool ($R_{\theta, k}$)		
$k \backslash \theta$	0	1
k_L	R_{0, k_L}	R_{1, k_L}
k_H	R_{0, k_H}	R_{1, k_H}

Fully pool (R)		
$k \backslash \theta$	0	1
k_L	R	R
k_H	R	R

Cross-pool ($R_a \neq R_b$)		
$k \backslash \theta$	0	1
k_L	R_a	R_b
k_H	R_b	R_a

Figure 3.1: Graphical definitions of *Semi-pool*, *Fully pool*, and *Cross-pool* for $R_{\theta, k}$. Identical colours denote equality. For precision, replace R with α .

To arrive to equilibria candidates, pairing $R_{\theta, k}$ and corresponding $\alpha_{\theta, k}$ is the next step. At first instance, one may believe that there are 3×3 cases. However, we can further eliminate the impossible cases given the nature of advisor's problem. To see this, we first outline all potential cases. Given unverifiable precision, the feasible combinations are listed in Table 3.1.

Table 3.1: Potential combinations of fee and precision forms

	Semi-pool $\alpha_{\theta,k}$	Fully pool $\alpha_{\theta,k}$	Cross-pool $\alpha_{\theta,k}$
Semi-pool fee $R_{\theta,k}$	✓		
Fully pool fee $R_{\theta,k}$	✓	✓	✓
Cross-pool fee			✓

Then, set up advisor's problem fully: In time 1, advisor determines $(\alpha_{\theta,k}, R_{\theta,k}(\alpha_{\theta,k}))$ to maximise

$$EU_A(\alpha_{\theta,k}, \alpha_{\theta,k'}, \hat{\theta}, \hat{k}; \theta, k) = R_{\theta,k} + k\alpha_{\theta,k}\pi$$

subject to incentive compatibility (IC) and individual rationality (IR) constraints:

Advisor's IC: $EU_A(\alpha_{\theta,k}^*, \alpha_{\theta,k'}^*, \theta, k; \hat{\theta}, \hat{k}) \geq EU_A(\alpha_{\theta,k}, \alpha_{\theta,k'}, \theta, k; \hat{\theta}, \hat{k})$ for all $(\alpha_{\theta,k}; \alpha_{\theta,k'})$ and for any (θ, k) .

Advisor's IR: $EU_A \geq 0$ (advisor has no other outside options).

Client's IR: Trading profit under rational expectations $\pi \geq R_{\theta,k}$.

To simplify the problem, I then work on the client's IR constraint:

Lemma 1. *Client's IR is binding: Trading profit under rational expectations $\pi = R_{\theta,k}$. Proof: see Appendix A.*

The problem then simplifies to finding the equilibrium expected trading profit, as the equilibrium level of fee $R_{\theta,k}$ given structure of $\alpha_{\theta,k}$. It then becomes clear that one signal structure might not realise all of the three fee structures (semi-pool, fully pool and cross-pool). In fact, plug in a given signal structure and compute expected trading profit, we end up with five equilibrium candidates. The first one being semi-pool fee and semi-pool signals:

$$\alpha_{k_L} \neq \alpha_{k_H} \quad R_k = (2\alpha_k - 1)\varepsilon(1 - \varepsilon)$$

The second is fully pool fee and fully pool signals:

$$\alpha_{\theta,k} = \alpha \quad \forall(\theta, k) \quad R_{\theta,k} = (2\alpha - 1)\varepsilon(1 - \varepsilon)$$

These are easy to spot. Other equilibria include a fully pool fee with semi-pool signals equilibrium:

$$\alpha_{k_L} \neq \alpha_{k_H} \quad R_k = \mathbb{E}_k((2\alpha_k - 1)\varepsilon(1 - \varepsilon)) \quad \forall(\theta, k)$$

; fully pool fee with cross-pool signals:

$$\alpha_{1,k_H} = \alpha_{0,k_L} = \alpha_0 \neq \alpha_1 = \alpha_{1,k_L} = \alpha_{0,k_H}$$

When $\mathbb{P}(k_H) = \mathbb{P}(k_L)$, this is close to case 3:

$$R_k = \mathbb{E}_k((\alpha_0 + \alpha_1 - 1)\varepsilon(1 - \varepsilon)) \quad \forall(\theta, k)$$

And finally, a cross-pool fee, cross-pool signals equilibrium:

$$\alpha_{1,k_H} = \alpha_{0,k_L} = \alpha_0 \neq \alpha_1 = \alpha_{1,k_L} = \alpha_{0,k_H}; \quad R_{\theta,k} = \frac{1}{4}(2\alpha_{\theta,k} - 1)$$

These equilibria feature the advisor's ability to determine their signal quality in a way that does not reveal the conflict through fee, but secretly conditional on the conflict scale (advisor's type) information she knows at the time of sending information to the client. This scenario is possible when signal performance, as a probability measure, is not contractable and unverifiable to client, highlighting the issue of moral hazard when advisor has a conflict of interest. Had this information be observable to client, those equilibria would not exist. ³

It is straightforward to verify that case 5 is not incentive compatible with advisor's types because the equilibrium level of precision (α_0, α_1) needs to satisfy both IC for k, k' . Given that in case 5, $R_{\theta,k}$ requires only one input in the pair (α_0, α_1) , the system of equations that solves (α_0, α_1) would have four equations for two unknowns. This overdetermined system cannot be solved. The logic expands to cases where $\mathbb{P}(k_H) \neq \mathbb{P}(k_L)$.

³When relaxing the assumption on market maker to the case where market maker does not know $R_{\theta,k}$, all fully pool equilibria with fees survive this robustness check.

Table 3.2: Equilibrium cases carried forward

Fee form $R_{\theta,k}$	Precision form $\alpha_{\theta,k}$	Notes
Semi-pool	Semi-pool ($\alpha_{k_L} \neq \alpha_{k_H}$)	Conflict-revealing fees.
Fully pool (R)	Fully pool (α)	Conflict-hiding; one common precision.
Fully pool (R)	Semi-pool ($\alpha_{k_L} \neq \alpha_{k_H}$)	Conflict-hiding; type-dependent precision.
Fully pool (R)	Cross-pool $\alpha_{1,k_H} = \alpha_{0,k_L}, \alpha_{1,k_L} = \alpha_{0,k_H}$	Cross-signalling precision.

Table 3.3: Admissible combinations of fee and precision forms

	Semi-pool $\alpha_{\theta,k}$	Fully pool $\alpha_{\theta,k}$	Cross-pool $\alpha_{\theta,k}$
Semi-pool fee $R_{\theta,k}$	✓		
Fully pool fee $R_{\theta,k}$	✓	✓	✓
Cross-pool fee			

There are four remaining cases, cases 1 to 4, to solve for analytical solutions, which is the centre for next section. Some comments to those equilibria are in place. Among these candidates, case 1 is *conflict-revealling*, that means observing $R_{\theta,k}$ allows inference on advisor's private value state k to full adjustment in beliefs. All other equilibria are *conflict-hiding*, because belief on k after observing $R_{\theta,k}$ stays at prior. The four cases carried forward in the analysis are listed in Table 3.2 (row: fee form; second column: precision form). To see what have changed after we solve the advisor's constraints, we have 3.3 rather than 3.1 before. For the rest of the paper I will use those terminologies as in line with literature. The key question to ask is whether the suspicion on conflict-hiding cases result in lower information efficiency, or higher advisor gains are indeed true. Another important question to resolve for regulators is, if there are cases for which advisor hides information on asset state, how can a regulator steps in to improve efficiency outcome?

3.2 A simple benchmark: no private incentives

To analyse the impact of private incentive, first look at a simple case with $k = 0$. Under this specific case of aligned incentive, as long as the profit function increases with precision levels, the advisor optimises with perfect information transmission. It is indeed the case for all potential equilibria, that the expected profit function π

(which is the fee generated, and so the expected payoff for the advisor at $k = 0$) increases with precision $\alpha_{\theta,k}$, so the below lemma would be true:

Lemma 2. *The baseline solutions (i.e. $k = 0$) to all cases would be setting $\alpha_{\theta,k} = 1$ with or without market noise.*

Chapter 4

Solving the Baseline Model

4.1 Baseline, case 1: semi-pool fee and semi-pool precision

Recall the advisor's problem

$$\max_{\alpha_k} EU_A = \underbrace{\varepsilon(1-\varepsilon)(2\alpha_k-1)}_{R_k} (1+k\alpha_k) \text{ s.t. } \text{Advisor's IC, IR}$$

Taking first-order conditions with respect to precision α_k , we obtain

$$\frac{dEU_A}{d\alpha_k} = \varepsilon(1-\varepsilon)((1+k\alpha_k)*2+k(2\alpha_k-1)) = \varepsilon(1-\varepsilon)(4\alpha_k k - k + 2)$$

The second-order condition holds for $k < 0$, where we obtain an interior solution. Then, the equilibrium precision is either an interior solution or two boundary solutions, depending on whether $\alpha_k^* = \frac{1}{4} - \frac{1}{2k} \in (\frac{1}{2}, 1)$ is true:

$$\alpha_k^* \in \left\{ \frac{1}{4} - \frac{1}{2k}, \frac{1}{2}, 1 \right\}$$

To characterise the set of parameters for equilibrium existence, first, solve the respective domains for the above three solutions; then check whether constraints are satisfied. To achieve the interior solution $\alpha_k^* \in (\frac{1}{2}, 1)$, $k \in (-2, -\frac{2}{3})$. Above $-\frac{2}{3}$, the objective increases in α_k and so $\alpha_k^* = 1$ and below -2 , the objective decreases in α_k and so $\alpha_k^* = \frac{1}{2}$. Then, we pin down the advisor's utility according to equilibrium

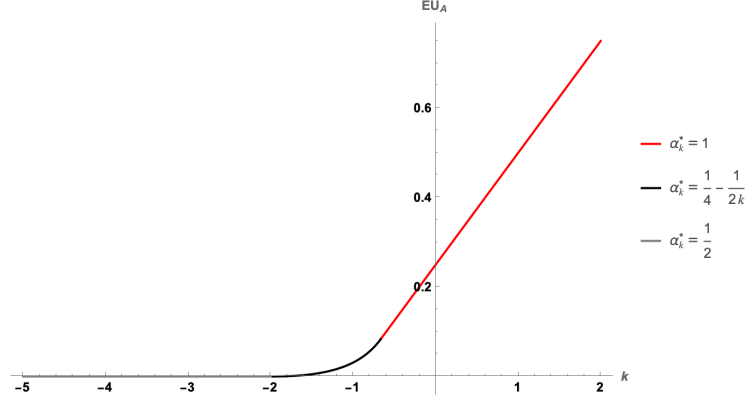


Figure 4.1: Conflict-revealing Equilibria: Expected payoff as a function of conflict k

solutions to see whether the constraints are satisfied:

$$EU_A^*(k) = \begin{cases} EU_A(\alpha_k^* = 1) = \varepsilon(1 - \varepsilon)(1 + k) & \text{if } k \geq -\frac{2}{3} \\ EU_A(\alpha_k^* = \frac{1}{4} - \frac{1}{2k}) = \frac{\varepsilon(1 - \varepsilon)(2 + k)^2}{8k} & k \in (-2, -\frac{2}{3}) \\ EU_A(\alpha_k^* = \frac{1}{2}) = 0 & \text{otherwise} \end{cases}$$

The IR constraint is satisfied as $EU_A \geq 0$ for these equilibrium level of expected payoffs. The advisor's IC constraint is satisfied by checking incentive compatibility of type k to type $k' \neq k$ solutions. From the graph 4.1, we can see that if $\alpha_{k'}^*$ is any of the boundary solutions, the IC constraints are satisfied. For the scenario with two interior solutions, the utility of type k mimic type k' is

$$EU_A(k, \alpha_{k'}^*) = \frac{\varepsilon(1 - \varepsilon)(k(k' - 2) + 4k')}{8(k')^2}$$

and its then easy to check for $k \in (-2, -\frac{2}{3})$ and $k' \in (-2, -\frac{2}{3})$, $EU_A(k, \alpha_{k'}^*) < EU_A(k, \alpha_k^*)$. Then, the above equilibrium characterisation satisfies all constraints.

Finally, I want to explain the cutoff of k for each subcase (in different colour). It is a result of checking that whether the equilibrium fee is indeed signalling the private state k , i.e. $R_{k_L} \neq R_{k_H}$. As $R_{\theta, k}$ is a linear function of precision $\alpha_{\theta, k}$. To get $R_{k_L} \neq R_{k_H}$, we need $\alpha_{k_L} \neq \alpha_{k_H}$. This statement can be transformed to such: the set of parameters (k_L, k_H) cannot fall in the boundary $\{\frac{1}{2}, 1\}$ at the same time. That gives the set of parameters $k_H \geq -\frac{2}{3}, k_L < -\frac{2}{3}$ or $k_H > -2, k_L < -2$. From

graphical representation in fig. 4.1, conflict-revealling equilibria only exists in three scenarios: $EU_A(k_L)$ in grey region, but $EU_A(k_H)$ is not; $EU_A(k_H)$ in red region, but $EU_A(k_L)$ is not; both $EU_A(k_L)$ and $EU_A(k_H)$ are in black region with distinct interior solutions.

This baseline scenario of conflict-revealling equilibria are in line with some public knowledge on good or bad advisors. The larger their conflict of interest, the (weakly) less information they will provide and lower profit they can achieve from the information market. In the model, we can find some support. The first case suggests that low type drops out of the information market because the conflict is too large for them to provide information, and high type exists in the market. There is a lower bound for high type to provide advice to client, i.e. the conflict is not too large for the high type, as expected. The second case means the high type advisor has sufficient incentive to always reveal the true asset value, but the low type is not. This reflects conventional wisdom: providing adequate incentives to advisor is necessary so that they are giving good advice to client. The interior solutions also presents a monotonic increasing relationship between α_k^* and k that mirrors this argument.

The model also translated to the following intuitive statement: if large enough conflict exists and advisor holds a certain extent of certainty about it, advisor signals direction and degree of direct conflict through the service and fee level they provide. By such construction we have in this section, the fee is a one-to-one mapping to the client's realised precision and client infers the conflict, which is often desirable from a transparency point of view. Practically, this argument corresponds to wide range of advisor qualities one might believe in the advisory market. We should note that this is not the only behaviour from advisors in the real world. Whether the advisor wants to voluntarily disclose their conflict at all time and use the fee as a signalling device is another question that would be answered later.

4.2 Baseline, case 2: fully pool fee and fully pool precision

Given fully pool signal structure, the advisor's problem is:

$$\max_{\alpha} EU_A = \underbrace{\varepsilon(1-\varepsilon)(2\alpha-1)}_R (1+k\alpha) \text{ s.t. } \text{Advisor's IC, IR}$$

As this problem is analogous to baseline case 1 except for the subscript k , immediately we know the below should be true. First, there exist a solution that both types reveal the truth:

$$\alpha_{k_L} = \alpha_{k_H} = 1 \quad \text{if } -\frac{2}{3} \leq k_L < k_H$$

Second, there exist a solution that both types provide no information:

$$\alpha_{k_L} = \alpha_{k_H} = \frac{1}{2} \quad \text{if } k_L < k_H \leq -2$$

For the remaining parameter combinations, where there exists a unique maximisation solution for case 1 that satisfies IC constraint, intuitively there is no possibility for a case 2 solution. From the previous section, maximisation problem yields the interior solution $\alpha_k^* = \frac{1}{4} - \frac{1}{2k}$. However, under the domain $k \in (-2, -\frac{2}{3})$, type k' does not want to pool with k due to violation of incentive compatibility constraints.

Then, the only fully pool equilibria comprise of either full truthtelling from advisors with sufficient incentive, or no information provided in the market as advisors has too much conflicts in the market. Graphically speaking, both private incentive parameters k_L and k_H lies in regions red or grey, respectively. Although conflicts are hided, the client has no risk to be deceived, as it is of advisor's benefit to provide true information to them. What about other types of conflict-hiding equilibria?

4.3 Baseline, case 3/4: fully pool fee and semi-pool/cross-pool precision

As said in the final part of section 1.2.1, at $\mathbb{P}(k_H) = \mathbb{P}(k_L)$ case 3 and case 4 are very similar. In those two equilibria, due to the nature of the fee function, aggregating the precision pairs $(\alpha_{\theta,k}, \alpha_{\theta,k'})$ set up by advisors would result in the same outcome. Therefore, I combine the two cases in this section and first analyse this simple case. The advisor's problem becomes:

$$\max_{\alpha_{\theta,k}, \alpha_{\theta,k'}} EU_A = \underbrace{\varepsilon(1 - \varepsilon)(\alpha_{\theta,k} + \alpha_{\theta,k'} - 1)}_R (1 + k\alpha_{\theta,k}) \quad s.t. \quad \text{Advisor's IC, IR}$$

To see how equilibrium can be solved, and the fact that there is no equilibrium with the given structure, consider the following logic. The advisor, places positive probability only on realised state k recoups the expected trading profit from precisions $\alpha_{\theta,k_L}, \alpha_{\theta,k_H}$ from client, who places positive probability to both states (k, k') . For any k , the decision rule for precision, α_k^* given choice of $\alpha_{\theta,k'}^*$ is

$$\alpha_{\theta,k}^*(\alpha_{\theta,k'}^*, k, k') = \arg \max_{\alpha_{\theta,k}} EU_A(\alpha_{\theta,k}, \alpha_{\theta,k'}, k) \quad s.t. \quad \text{Advisor's IC, IR}$$

EU_A increases in $\alpha_{\theta,k'}$, shown by positive FOC of $\alpha_{\theta,k'}$, $\frac{dEU_A(k)}{d\alpha_{\theta,k'}} = 2\mathbb{P}(k')\varepsilon(1 - \varepsilon) > 0$. Then, type k advisor promises $\alpha_{\theta,k'}^* = 1$ and recoups half of maximum expected trading profit $\pi(\alpha = 1)$ as associated fee. This is true regardless of high or low type or even type probability. However, the choice is either not incentive compatible for type k' : $EU_A(k')$ is not always increasing in $\alpha_{\theta,k'}$, or when it does, the solution degenerates to the fully pool solution, $\alpha_{\theta,k}^* = \alpha_{\theta,k'}^* = 1$. Alternatively, the first scenario also violates client's participation constraint: once a proposed fee $R_{\theta,k}$ satisfies $R_{\theta,k} \in (\frac{1}{2}\pi(\alpha_{\theta,k} = 1, \alpha_{\theta,k'} = 1), \pi(\alpha_{\theta,k} = 1, \alpha_{\theta,k'} = 1))$, the client would know that type k is not the truthtelling type, and adjusted fully the belief, and realised that the charged fee would be lower than the expected trading profit given the adjusted belief. So, the equilibrium with conflict-hiding and partial truthtelling would not exist. In the case that advisor might want to charge $R_{\theta,k} < \pi$, we can

use a similar reasoning as in lemma 1 and prove by contradiction, assuming first that there is indeed a optimal level of fee satisfying this criterion. The reasoning uses three elements. First, the fee increases in precision. Second, the advisor cares only about realised precision for state k . Third, the client cares about their expected precision. As type probability is not a factor in the reasoning, the same rationale should extend to cases with $\mathbb{P}(k_H) \neq \mathbb{P}(k_L)$ and applies to both cases 3 and 4.

This has important implications on equilibrium existence. First, one observes the FOC of $\alpha_{\theta,k'}$ resulted from the fact that advisor places zero weight on $\alpha_{k'}$ in expected private value (relative to π). However, this is not true for $R_{\theta,k}$, as value of information in a pooling fee equilibrium would contain information from both private incentive realisations, i.e. $(\alpha_{\theta,k}, \alpha_{\theta,k'})$. This asymmetry directly lead to no enough uncertainty for the advisor to hide type while credibly send information in precision indicated by the value of fee $R_{\theta,k}$. As one can interpret the current specification for private value as a *direct* measure for conflict of interest, measuring signal performance, one might also ask the effect of the incentive function specification Q have to the model results. From the above logic, the advisor's maximisation problem suggests that, in general, when marginal increase of utility EU_A through $\alpha_{\theta,k}$ are not equal across types (k, k') , or, not guaranteed as positive or negative, the problem persists. The former criterion suggests there is an incentive-compatible interior solution for $\alpha_{\theta,k}$, which means advisor is at least not truthtelling at a node k . The latter criterion suggests that the solution $\{\frac{1}{2}, 1\}$ is incentive compatible, which means advisor at half time tells the truth, and half time babbling. These two criteria sees through private values that creating conflict-hiding equilibrium with advisor hiding information on asset value. So, the above no-equilibrium intuition is not limited to a specific incentive function. With specifications similar to the baseline, that is, advisor payoff asymmetries between weighting on precision in the fee charged ($R_{\theta,k}$) and the private value function, we can extend the intuition described here.

One might argue that the baseline is an oversimplification to the reality. The real world scenario might feature advisor's private values that are not directly associated with signal performance. Those *indirect* conflict of interest can involve

commissions that contingent on market prices (e.g. advisor charges a proportion of the transaction price as execution cost), or other related service lines. To tackle the situation, regulators classify advisors and brokers into categories according to their functionalities in the market, and monitor other entities that has advisory services on their alternative service lines, imposing restrictions (e.g. Volcker rule) to ensure that conflicts does not harm their advisees. Another widely discussed policy consists fee unbundling in advisory services, for which EU regulations restrict the form of research and transaction costs of an advisor, and charging the two costs in a combined fee is no longer allowed. These procedures and checking the above criteria surely help with understanding whether the bad conflict-hiding equilibria can happen. However, from the model, we might notice that the true problem lies in agents hiding the private incentive state, and client can only infer the private state at prior belief. So, if we allow mandatory disclosure of the parameter k truthfully, client can fully adjust the belief, forcing any existing information provider to charge fee as in case 1 and a conflict-revealling situation arise. In practice, this would correspond to agent's relative scale of conflict of interest with the client seeking for advice. For any case where advisor knows exactly the degree of such conflict of interest, the rule should eliminate any conflict-hiding equilibria without truthtelling, as the only case would remain as fully pool with full information precision (case 2).

However, why mandatory disclosure is not the fully effective solution in the real world? Why advisor would not value transparency themselves so disclosure rules might not be necessary at all? We can see real examples that disclosure is not fully reflecting the degree and/or direction of conflict (e.g. an accounting fraud). Another case might be that k is uncertain when advisor charges the fee, so they might not be able to disclose it before the client makes the purchase decision. Abundant real-world motivations exist as uncertainty in a firm's conflict of interest, which is not covered in the baseline model. For example, banks advising the client might face short-selling constraints when they engage in asset market activities, such as market making. Another example might be that an advisor is advising another client with conflict of interest to the first client, however, this is not known when the first

client matches with the advisor.

To show those cases, I consider a modified model below and illustrate how case 3 or 4 can emerge from advisor's uncertainty for future conflicts.

Chapter 5

Advisor Type Uncertainty: A Modified Problem

To convey what I have mentioned above, i.e., advisor cannot reveal conflict due to uncertainty, and so has no choice but to pool fee, I make changes to the model. From a timeline perspective, it meant we have changed the advisor's information set at time 1, from $F_A = \{\theta, k\}$ to $F_A = \{\theta\}$; and at the start of time 2, $F_A = \{\theta, k\}$. This means advisor has the choice of semi-pool or cross-pool signal precisions at time 2, but can only present a single fee level at time 1. To represent the general case for advisor's problem

$$\max_{\alpha_{\theta,k}, \alpha_{\theta,k'}} EU_A = \underbrace{\varepsilon(1-\varepsilon)(\alpha_{\theta,k} + \alpha_{\theta,k'} - 1)}_R \left(1 + \frac{1}{2} \sum_k k \alpha_{\theta,k}\right) \quad s.t. \quad \text{Advisor's IC, IR}$$

From the private incentive construction, we can see that the advisor now cares about the expected private payoff reflected in $\frac{1}{2} \sum_k k \alpha_{\theta,k}$, i.e. $E(k \cdot \mathbb{1}(Y = \theta) | \theta)$, which is what we described.

To check advisor's IC, $EU_A(\alpha_k^*, \alpha_{k'}^*) \geq EU_A(\alpha_k, \alpha_{k'})$, two cases for deviations should be noted. One is for the small deviations for type k , given the other type plays $\alpha^*(k, k')$; this is checked by FOC. The other is large deviations, which means the low type to no information, and the high type to truthtelling. This condition I check manually and result in some finite number of inequalities restricting equilibria existence.

Table 5.1: Admissible combinations of fee and precision forms

	Semi-pool $\alpha_{\theta,k}$	Fully pool $\alpha_{\theta,k}$	Cross-pool $\alpha_{\theta,k}$
Semi-pool fee $R_{\theta,k}$			
Fully pool fee $R_{\theta,k}$	✓	✓	✓
Cross-pool fee			

Candidate equilibria are cases 2-4 corresponding to characterisation in baseline (section 4). Again, the cross fee solution is also not incentive compatible. Table 5.1 shows a graphical representation.

5.1 Type uncertainty, case 2: fully pool fee and fully pool precision

The problem is a variation of baseline case 1 in the baseline, because for the modified problem, $\alpha_{\theta,k} = \alpha_{\theta,k'}$ implies that

$$EU_A = \underbrace{\varepsilon(1-\varepsilon)(2\alpha-1)}_R (1 + \mathbb{E}(k)\alpha)$$

and note that we are replacing k in baseline's utility function to $\mathbb{E}(k)$. As $\mathbb{E}(k)$ is but a constant, arguments in the baseline model remains to hold. Then,

$$EU_A^*(k) = \begin{cases} EU_A(\alpha = 1) & \text{if } \mathbb{E}(k) \geq -\frac{2}{3} \\ EU_A(\alpha^* = \frac{1}{4} - \frac{1}{2\mathbb{E}(k)}) & \mathbb{E}(k) \in (-2, -\frac{2}{3}) \\ EU_A(\alpha = \frac{1}{2}) & \text{otherwise} \end{cases}$$

The below shows arguments to check large deviation.

If $\mathbb{E}(k) \leq -2$, given the belief, need to check whether high type can benefit from providing more information (deviate to case 3) and the condition must hold even for highest action $\alpha_{k_H} = 1$. This result in additional constraints on k_H :

$$k_L < k_H \leq -\frac{2}{3}(2 + k_L)$$

If $\mathbb{E}(k) \geq -\frac{2}{3}$, need to check whether the low type wants to provide less information, and the condition must hold even for lowest action $\alpha_{k_L} = \frac{1}{2}$. The additional constraints on k_L are:

$$\begin{cases} \frac{2}{k_L} \left[\frac{2}{k_L} - \frac{1}{2}k_H - 1 \right] > 1 & \text{if } \mathbb{E}(k) \geq -\frac{2}{3}, k_H > k_L > 0 \\ \frac{2}{k_L} \left[\frac{2}{k_L} - \frac{1}{2}k_H - 1 \right] < \frac{1}{2} & \text{if } \mathbb{E}(k) \geq -\frac{2}{3}, k_H > 0, k_L < 0 \end{cases}$$

If $\mathbb{E}(k) \in (-2, -\frac{2}{3})$, we want to check both low type to no information and high type to truthtelling. The relevant conditions are already defined above, except that we need to replace the range for $\mathbb{E}(k)$. Then, unlike in the baseline, there are some fully pool equilibrium with interior solutions $\alpha^* = \frac{1}{4} - \frac{1}{2\mathbb{E}(k)}$.

5.2 Type Uncertainty, Case 3/4: Fully Pool Fee and Semi-Pool/Cross-Pool Precision

These conflict-hiding equilibria has the same feature: the fee does not reveal private incentive but the information precision corresponded to each private incentive realisation would be different. Due to unobserved and unverifiable information precision in the model, the advisor can determine the *realised* information precision after the fee is charged from the client. However, due to type uncertainty and the states (θ, k) being independent, the advisor would assign the same belief to the private state as the client. That is, $\mathbb{P}(k|\theta) = \mathbb{P}(k)$, and assign both precision parameters $(\alpha_{\theta,k}, \alpha_{\theta,k'})$ as realised value with respective k probabilities. This is the basis for equilibrium existence.

To see this is the case, recall the argument in the baseline model where there is no type uncertainty. As advisor assign zero probability to the precision parameter that is not realised, $\alpha_{\theta,k'}$, the advisor sets $\alpha_{\theta,k'} = 1$ which is either not incentive compatible to the advisor or violates the individual rationality constraint of the client. In the revised case, the same argument does not carry over totally. If the advisor sets $\alpha_{\theta,k'} = 1$, it is a result of the FOC $\frac{dEU_A}{d\alpha_{\theta,k'}} > 0$ for all $\alpha_{\theta,k'} \in [\frac{1}{2}, 1]$, for which $\frac{dEU_A}{d\alpha_{\theta,k'}}$ is a function of *both* $(\alpha_{\theta,k}, \alpha_{\theta,k'})$. Thus, the condition for FOC

becomes $\frac{\partial EU_A}{\partial \alpha_{\theta,k'}}(\alpha_{\theta,k}^*, \alpha_{\theta,k'}) > 0$ and is *jointly* determined by values $\alpha_{\theta,k'} \in [\frac{1}{2}, 1]$ and $\alpha_{\theta,k}^*$ determined by its respective FOC. We then need to figure out all possibilities, as this argument give rise to potential $\alpha_{\theta,k'} \neq 1$. Intuitively, we have the following: $(\alpha_{\theta,k} = 1, \alpha_{\theta,k'} = 1)$ which is a boundary fully pool equilibrium; $(\alpha_{\theta,k}^* \in (\frac{1}{2}, 1), \alpha_{\theta,k'} = 1)$; $(\alpha_{\theta,k} = \frac{1}{2}, \alpha_{\theta,k'} = 1)$; $(\alpha_{\theta,k} \in (\frac{1}{2}, 1), \alpha_{\theta,k'}^* \in (\frac{1}{2}, 1))$; $(\alpha_{\theta,k} = \frac{1}{2}, \alpha_{\theta,k'}^* \in (\frac{1}{2}, 1))$ and finally $(\alpha_{\theta,k} = \frac{1}{2}, \alpha_{\theta,k'} = \frac{1}{2})$ which is the other boundary fully pool equilibrium.

Formally, for each of the above case, use the respective criteria to test whether a given solution range can be possible, pairing with the advisor's IR constraint $EU_A > 0$:

1. If $\alpha_{\theta,k}^* = 1$: $\frac{\partial EU_A}{\partial \alpha_{\theta,k}}(\alpha_{\theta,k}, \alpha_{\theta,k'}^*) > 0$ for optimal $\alpha_{\theta,k'}$, $\alpha_{\theta,k'}^*$, and $\alpha_{\theta,k} \in [\frac{1}{2}, 1]$.
2. If $\alpha_{\theta,k}^* \in (\frac{1}{2}, 1)$: $\frac{\partial EU_A}{\partial \alpha_{\theta,k'}}(\alpha_{\theta,k}^*, \alpha_{\theta,k'}^*) = 0$ for optimal $\alpha_{\theta,k'}$, $\alpha_{\theta,k'}^*$.
3. If $\alpha_{\theta,k}^* = \frac{1}{2}$: $\frac{\partial EU_A}{\partial \alpha_{\theta,k}}(\alpha_{\theta,k}, \alpha_{\theta,k'}^*) < 0$ for optimal $\alpha_{\theta,k'}$, $\alpha_{\theta,k'}^*$, and $\alpha_{\theta,k} \in [\frac{1}{2}, 1]$.

and test the four cases. The exact procedure, as well as conditions for equilibrium to exist, are in Appendix C. Fig. 5.1 shows the graphical representation for equilibrium domains.

The conjecture is that $\alpha_{\theta,k_H} \neq 1$ is still of low possibility, as the fee $R_{\theta,k}$ is increasing in both precision parameters, and to have $\alpha_{\theta,k_H} < 1$ requires the private benefit to be sufficiently negative to induce this, i.e. k_H sufficiently small. The same applies to $k_L < k_H$. However, at least one of $\alpha_{\theta,k} > \frac{1}{2}$, which is necessary for a case 3 equilibrium to exist, $1 + \frac{1}{2} \sum_k k \alpha_{\theta,k} > 0$ need to hold for advisor's IR constraint. This poses a lower bound for both k_L and k_H . As k_H becomes larger, the benefit of setting $\alpha_{\theta,k_H} < 1$ decreases.

After checks (in Appendix C), we have the two candidate subcases to work on: $(\alpha_{\theta,k}^* \in (\frac{1}{2}, 1), \alpha_{\theta,k'} = 1)$ and $(\alpha_{\theta,k} = \frac{1}{2}, \alpha_{\theta,k'} = 1)$. I name the two cases case 3.1 (truthtelling high type and partial truthtelling low type) and case 3.2 (truthtelling high type and babbling low type) respectively. Due to the nature of the game, the good type would provide weakly more information. Then, the below notation observes this pattern: $k = k_L$ and $k' = k_H$.

5.2.1 Case 3.1: Truthtelling High Type and Partial Truthtelling Low Type

To solve for analytical solutions, the idea is to solve for α_{θ,k_L}^* such that plugging optimal $\alpha_{\theta,k_H}^* = 1$ and

$$\frac{\partial EU_A}{\partial \alpha_{\theta,k_L}}(\alpha_{\theta,k_L}^*, 1) = 0$$

Then, the solution is $\alpha_{\theta,k_L}^* = -\frac{2+k_H}{k_L}$. Incentive compatible solution requires

$$\frac{\partial EU_A}{\partial \alpha_{\theta,k_H}}(\alpha_{\theta,k_L}^*, \alpha_{\theta,k_H}) > 0$$

for $\alpha_{\theta,k_H} \in [\frac{1}{2}, 1]$ and

$$\alpha_{\theta,k_L}^* = \arg \max_{\alpha_{\theta,k_L}} EU_A$$

Additionally, we need $\alpha_{\theta,k_L}^* = -\frac{2+k_H}{k_L} \in (\frac{1}{2}, 1)$ to allow equilibrium existence.

In this case the equilibrium value of R is $\frac{1}{2} \left(\pi + \varepsilon(1 - \varepsilon) \left(-\frac{2(2+k_H)}{k_L} - 1 \right) \right)$, the sums represent the information content of full truth if advisor realises good type, and only partial information if realises the bad type. The expected value of the information from good type is $\frac{1}{2}\pi$, half of the maximum trading profit one can get.

5.2.2 Case 3.2: Truthtelling High Type and Babbling Low Type

Following a similar reasoning, first, to have $\alpha_{\theta,k_L}^* = \frac{1}{2}$, we need

$$\frac{\partial EU_A}{\partial \alpha_{\theta,k_L}}(\alpha_{\theta,k_L}, \alpha_{\theta,k_H}^* = 1) < 0$$

for $\alpha_{\theta,k_L} \in [\frac{1}{2}, 1]$. Also, $\alpha_{\theta,k_H}^* = 1$ requires

$$\frac{\partial EU_A}{\partial \alpha_{\theta,k_H}} \left(\alpha_{\theta,k_L}^* = \frac{1}{2}, \alpha_{\theta,k_H} \right) > 0$$

for $\alpha_{\theta,k_H} \in [\frac{1}{2}, 1]$.

In this case the equilibrium value of R is $\frac{1}{2}\pi$, representing the advisor either selling truth or tells nothing.

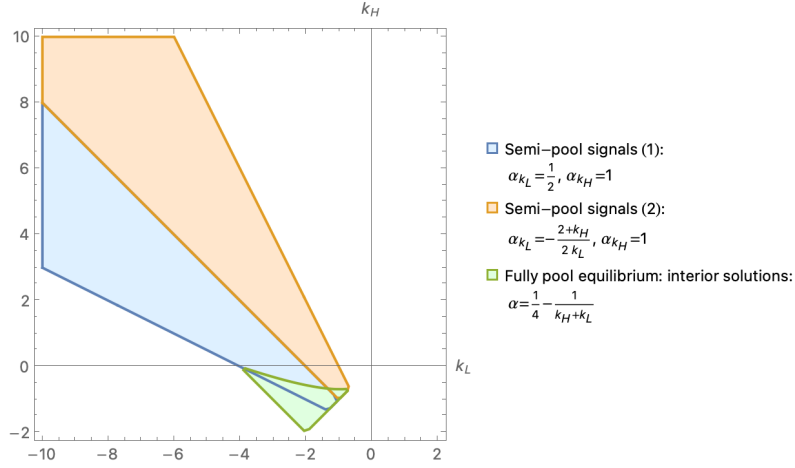


Figure 5.1: Equilibria with type uncertainties: Interior solutions

5.3 Understanding the Equilibria

This section aims to understand the outcomes from all three conflict-hiding equilibria in this type uncertainty world illustrate in Fig. 5.1. In particular, the equilibrium multiplicity and the trade off between advisor's private value and information provision generates interesting patterns. The most straightforward analysis lies in efficiency outcomes and comparative statics between efficiency and private incentive parameters k , as this determines the need and effectiveness of any regulations with regard to advisor's private benefit. Surprisingly, the information efficiency for the low type is nonmonotonic to private benefit scale k_L under some cases, and I will elaborate on its implications below. Moreover, a similar comparison between advisor's utility and information efficiency can be made to understand whether the traditional worries with moral hazard: advisor's improving utility at expense of the client. In this model, the statement appears to be partially, but not totally true, as with type uncertainty's multiple equilibria, the mechanism that drives towards higher advisor's utility and higher information efficiency can be the same *or* different channels, and this becomes clear when understand advisor's tradeoff. We can also infer the degree of necessity to resolve the advisor's uncertainty in conflict to improve efficiency outcomes. A useful comparison for the analysis is to work out, with given levels of (k_L, k_H) , equilibria in sections 2 or 3.

5.3.1 Summary of Analytical Results

Under the type uncertainty, the focus is on three cases with interior solutions. Fig. 5.1 shows a graph to visualise the equilibrium constraints for all these equilibria in the area satisfying $k_L < k_H$. Case 2, shaded green, includes all combinations with a fully pool equilibrium and $\alpha^* \in (\frac{1}{2}, 1)$. This green area lies in domains $k_H < 0$, with the relatively good type some degrees of conflict, though the magnitude is not very large. Case 3 equilibria features the high type selling fully informative signals, i.e. precision equal to 1. Case 3.1 in orange summarises the semi-pool equilibrium with the low type selling some informative signals. Next to case 3.1 is case 3.2, shaded blue, with the low type selling uninformative signals.

Those two subcases in case 3 occupies adjacent regions can be explained in algebra. When k_L decreases, given a value of k_H , the optimal value of $\alpha_{\theta, k_L}^* = -\frac{2+k_H}{k_L}$ decreases in k_L , and when $-\frac{2+k_H}{k_L}$ drops below $\frac{1}{2}$, it corresponds to utility EU_A decreases in $\alpha_{\theta, k_L} \in [\frac{1}{2}, 1]$ and thus the optimal $\alpha_{\theta, k_L}^* = \frac{1}{2}$. From intuition, it means when the low type's conflict rises ($k_L < 0$ and drops), within case 3 equilibrium, the contingent plan of advisor leaned to providing lower information content to the client, something easy to understand. Advisor faces a problem of uncertainty about being incentive aligned with the client, or not. If incentive is aligned, she wants to provide as much information as possible and so $\alpha_{\theta, k_H}^* = 1$; otherwise, the opposite applies. However, to credibly charge the fee and satisfy the incentive compatibility constraints, she optimally determines the low type information had k_L realised as the type after time 1, as a joint problem described in section 3.2.

Case 2 describes a different pattern in advisor's decision. It means advisor adjusts both decision rules at k_L and k_H together. Unlike in case 3, which only adjusts precision at k_L , case 2 exhibits one-to-one changes in such decision rules. Therefore, the optimal α^* in this setting becomes a function of both k_L, k_H , as a result of full pooling. Notice that this type of equilibrium can have three results, but other two cases, fully pool with truthtelling solution $\alpha^* = 1$ or babbling solution $\alpha^* = \frac{1}{2}$, are not contrasted with the interior solution equilibria. They are more comparable with the baseline solutions that yields the same solution, but the domains for

such solutions can be affected by type uncertainty, which would be analysed later. For the interior solutions in green region, $\alpha^* = \frac{1}{4} - \frac{1}{2\mathbb{E}(k)}$, and only located in some moderately negative k_H , i.e. even the good type advisor has some limited degree of conflicts. The low type has even more conflict. Intuitively, it creates incentive for the high type to pool with low type when both types present some conflict; the magnitude is determined by type uncertainty, but the direction is qualitatively the same. By such pooling, one would expect from the beginning that low type has higher information content produced in case 2 than case 3, at expense of high type pooling with them. However, the overall information efficiency becomes hard to conjecture on without further analysis.

It should be noted that the result depends on the form of private benefit being signal performance, as justified in the model setup. The exact analytical solution of this model also depends on the prior probabilities $\mathbb{P}(\theta)$ and $\mathbb{P}(k)$. However, even without any complications or asymmetric prior in the state space, the model realises some unanticipated results in equilibrium precision levels.

5.3.2 Substitutability of Advisor's Income Sources

To build on the results from the model, more generally, advisor balances two income sources under type uncertainty: aggregate efficiency ($\alpha_{k_L} + \alpha_{k_H}$) where higher value lead to higher fee; and private value ($k_L\alpha_{k_L} + k_H\alpha_{k_H}$) which changes to a weighted sum of precision compared to baseline. Immediately, when $k > 0$, the two factors exhibit *complementarity* at node k , and when $k < 0$, the two factors exhibit *substitutability* at node k .

Applying this argument to case 3, with $k_H > 0$ and $k_L < 0$, we have the following. $k_H > 0$ means that complementarity between the two factors at $k = k_H$ leads to $\alpha_H^* = 1$; and $k_L < 0$ means that substitutability between the two factors at $k = k_L$ leads to $\alpha_L^* \in [\frac{1}{2}, 1)$. This is very similar to the intuition from direct observation to the analytical results. Also, with $k_H > 0$, there is no incentive for the advisor to pool both types in precision because there are always benefit for the high type to tell the truth and improving private benefit, and the low type to provide less information because of the same channel, and so there are no case 2 (the

green region) with interior solution under such parameter combinations. From this mechanism, it is also evident that low type precision monotonically decreases with the corresponding conflict scale k_L , because the substitution between information precision and private value is linear.

However, when $k_H < 0$ the problem becomes harder to dissect and the monotonicity result collapses for α_{θ, k_L}^* . From the previous argument, advisor's fee, represented by aggregate precision, and private values substitutes each other. However, both case 3 and case 2 provide ways to substitute the income. Case 3 offers to reduce aggregate precision through only k_L but not k_H , and case 2 reduces the aggregate precision at both states. Intuitively, this deals with the relative strength of substitutability between the two types. Case 3 decreases k_L more drastically than in case 2. Due to this sensitivity difference, if $|k_H|$ is relatively small compare to $|k_L|$, advisor can find it more beneficial to provide much less information at k_L , at the cost of sacrificing private value at k_H by setting $k_H = 1$. However, there are some constraints for those equilibria, one of them being that $\alpha_{\theta, k_L}^* \in [\frac{1}{2}, 1)$. Case 3.1 happens to be more sensitive to changes in k than case 2, so this faster decay to lower values would drive out $\alpha_{\theta, k_L}^* = -\frac{2+k_H}{k_L}$ to $\frac{1}{2}$ much faster than case 2. Then, at a point, when the advisor moves from case 3.1 (orange), to case 3.2 (blue). However, the precisions for low and high types $\{\frac{1}{2}, 1\}$ also need to satisfy the advisor's participation constraints, and this is determined by the sensitivity of private value to k_L . However, by taking derivative of private incentive in case 3.2 and case 2, the latter also shows a lower sensitivity to k_L ($= \frac{1}{4}$) than that of case 3.2 ($= \frac{1}{2}$). And so, when k_L finally drops out of the case 3 domains due to utility EU_A decays faster to below zero, case 2 with interior solution (green) still survives, and moving from $\alpha_{\theta, k_L}^* = \frac{1}{2}$ to $\frac{1}{4} - \frac{1}{2\mathbb{E}(k)} = \frac{1}{4} - \frac{1}{k_H + k_L} \in (\frac{1}{2}, 1)$ suggests that α_{θ, k_L}^* can be non-monotonic to k_L for some $k_H < 0$ and $|k_H|$ below some threshold: a discontinuity between information provision and the private benefit scale for a potential low type. A graphical illustration is shown in fig. 5.2 with $k_H = -\frac{1}{2}$.

This highlights the regulatory difficulties when restraining advisor's conflicts of interest under conflict uncertainty. The model takes rather abstract way of mod-

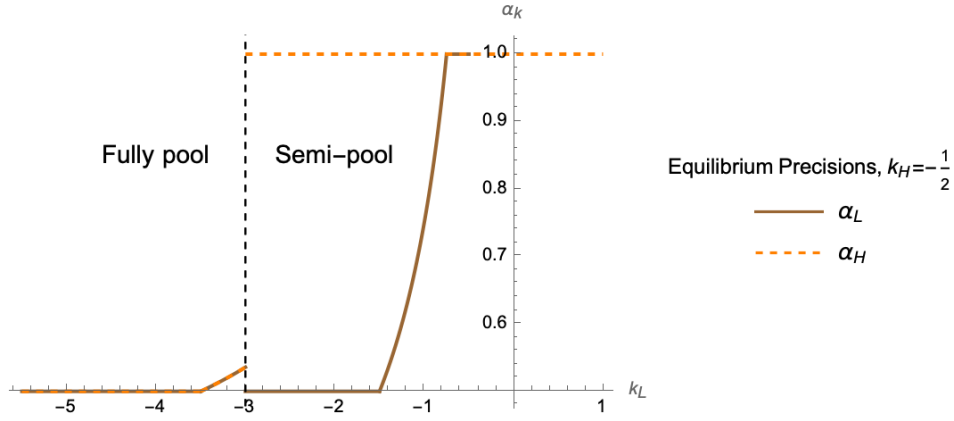


Figure 5.2: An example with $k_H = -\frac{1}{2}$. Dashed line: α_{k_H} . Brown line: α_{k_L} .

elling the conflict as a scale to the interaction term between the advisory payment and ex post signal precision, which captures the performance of advisor's signals. In practice, those conflicts can relate to advisor's commissions and kickbacks, representing conflicting fees; or, it can represent advisor's reputation costs of providing wrong signals. A natural argument is to restrict those conflicts if there exists one. In effect, an advisor's unknown incentive and unverifiable information precision allows advisor to shift between equilibria, rather than improving precision for information precision within one equilibrium, which obeys a single decision rule. Also, although different with cheap-talk setting of communication games, conflict-hiding equilibria with type uncertainty in the model generates similar patterns of triggering more information. Also, it underlines the importance of providing fully aligned incentives for the good type under type uncertainty, because if $k_H \geq 0$ these equilibrium multiplicity and nonmonotonic results vanish, and end up with the traditional stance that restricting conflict gives rise to better information production. These results generated from equilibrium multiplicity and implications from strategic complementarity/substitutability results draws similar patterns with Dow et al. (2017), with multiple equilibria when firm producing information on an asset (the firm's value itself), and the factor driving the multiple equilibria are also ex post variable realisations. However, Dow et al. (2017) achieve it through learning in asset market and endogenous decision of ex post outcome, where in this model, learning happens in both the information market on private value state, which its

realisation is exogenously determined, and the asset market, where market maker's prices reflect rational expectations to precisions, which is an endogenous plan by the advisor, but realisations are determined also exogenously.

5.3.3 Equilibria Comparisons

We also want to know whether there are jumps in those equilibrium multiplicity region. I adopt a Pareto dominance selection technique, that advisor plays the strategy with the highest utility gain among those equilibria, and client has rational expectations to the advisor's strategies. By assessing this criterion and comparing the aggregate efficiency levels between equilibria, we can also answer the question under conflict uncertainty scenario, that whether under the semi-pool equilibrium, the advisor is secretly contingent actions on unrealised private incentive at time of the sale, to gain higher utility but hurting the client. That is, whether lower aggregate efficiency pairs with a higher utility and its linkage to semi-pool equilibria. Calculating utility differences between cases 3.1 and 2 shows that the former gains higher utility at expense of lower efficiency, and utility gains are through lower private value losses at $k = k_L$. Then, under type uncertainty, it is possible for a semi-pool equilibrium to provide a higher expected value of information, while also gains a higher advisor utility, than a fully pool equilibrium that exists in the same parameter combination. Again, the reason of this result comes from the same trade-off between the two income sources and sensitivity to state k . Case 3.1 operates primarily through reducing private value losses at k_L , i.e. reducing information precision α_{θ, k_L}^* along k_L , and α_{θ, k_H}^* is not changed throughout. The equilibrium sacrifices aggregate precision through α_{θ, k_L}^* for lower private value losses at k_L , which becomes the dominant channel for utility changes.

Similar comparison operated to cases 3.2 and 2 shows an uncertain conclusion on utility and private value losses. However, the former has higher aggregate efficiency. Intuitively, the result would depend on two factors: level of $k_H + k_L$ which determines the aggregate efficiency and private value of case 2; and $\frac{1}{2}k_L + k_H$ which is the private value of case 3.2. However, the equilibrium constraints suggest that $k_H + k_L \leq -2$ which then implies $2 \left(\frac{1}{4} - \frac{1}{k_H + k_L} \right) \leq \frac{3}{2}$, which means there is a def-

inite ranking on aggregate efficiency: case 3.2 dominates. However, the private value comparison can at best restrain case 2 to $\frac{1}{4}(k_L + k_H) - 1 \leq -\frac{3}{2}$ and case 3.2 to $\frac{1}{2}k_L + k_H \geq -2$, but not exact ranking, and thus the utility comparison also ends with uncertain conclusion. One conjectured reason links to the advisor's relative substitutability between two income sources. The information precision rules of case 3.2 set $\alpha_{\theta, k_L}^* = \frac{1}{2}$ and $\alpha_{\theta, k_H}^* = 1$ and so the expected value of private benefit at k_L and k_H is weighted 1 : 2. That ratio is 1 : 1 for case 2. Then, which equilibrium leads to lower private value loss would depend on the *ratio* between k_L and k_H . The mechanism is entirely different with comparison between aggregate efficiency which is only restrained to the value of $k_H + k_L$, and as $k_H < 0$, the values are measured in absolute magnitude rather their relative values to another.

As a summary, under conflict uncertainty, whether advisor is using a semi-pool equilibrium to deliberately provide lower information content in their advice has no one-cut conclusions. When advisor's types presents not too much difference but the good advisor has a minor conflict, then advisor indeed is prone to the above worries. Though a fully pool equilibrium is a better precision outcome, it cannot be implemented using a verifiable method. Another form of regulation is to simply avoid such regions by enforcing the good type to be fully aligned. Otherwise, a semi-pool equilibrium with good type tells the truth would actually help with information production, even if the bad advisor chooses to produce nothing. The advisor's ex ante expected utility is not necessarily linked to information precision in the latter case. However, this creates further jumps in equilibrium within case 3.2 (blue region) that overlaps with case 2 (green region) and discontinuity of information provided by the realised low type. (A figure is provided in Appendix, fig. C.1.) This unpredictability poses regulatory challenges with those parameter combinations, and in the case, it is not clear that whether ex post transparency would be effective, without a clear idea the type of equilibria one is comparing under the conflict scale. (Kartal and Tremewan (2018); Ismayilov and Potters (2013); Behnk et al. (2014))

Sections 4.3.2-4.3.3 identify the parameter combinations of private incentives that are problematic under type uncertainty. In short, providing sufficient incen-

tives to the good advisors are always important to eliminate undesired outcomes. Restricting low type advisor, on the other hand, might not be helpful due to discontinuities in information production by the realised low type, and the relative substitutability between income sources.

5.3.4 Resolve the Conflict Uncertainty or Not?

If the uncertainty is resolved, the advisor faces case 1 or 2 equilibria in the baseline, where she follows a different decision rule for α_{θ, k_L}^* as only function of the realised type k_L . In particular, I focus on two dimensions. First, whether removing the conflict uncertainty in this section and revert back the advisor's problem to that of the baseline model can improve information precision. Second, how can regulations on private incentive scales k combine with relieving type uncertainty to enhance advisor's information provision.

First, notice that conflict uncertainty can move a good type advisor from full to partial truth-telling. The baseline case only emphasises the importance of *relative* incentive alignment, which means k is above some threshold $\bar{k}$ not necessarily equal to zero; and $\bar{k} = -\frac{2}{3} < 0$ in the current model specification. As long as the realised conflict is not sufficiently large, an advisor can tell the truth in the no-conflict-uncertainty case, because the objective for good type remains increasing within the domains, and thus the optimal solution $\alpha_{k_H}^* = 1$. The same statement is not true with universal case in conflict-uncertainty case, and additional constraints are made on k_L , i.e., when k_L is greater than a threshold value and $\bar{k} < k_H < 0$. For example, in fig. 5.2, I have shown that when k_L is low enough, it can drive k_H to also pool with k_L . With uncertainty, an advisor always attach positive probability to being low type under this scenario, and thus the precision for high type can jointly depend on both k . However, in the no-conflict-uncertainty case, decision at k_H only depends on values of k_H . This dependency difference results in the above observations, if no uncertainty results in case 1 or 2 equilibria with high type provides full information, but with uncertainty equilibrium results in case 2, the fully pool equilibrium with type interior solutions.

However, conflict uncertainty can sometimes generate some more degree of

communication for the good type. For example, the baseline suggests when $k_H < \bar{k}$, the good advisor provides partial information. However, with type uncertainty, it is possible for the high type to perform truthtelling under the semi-pool equilibrium solution, as long as k_L falls within a moderate region (orange and blue regions in fig. 5.1). In this case, the difference of conflict between the high and low type is moderate, and conflict of the high type does not give full incentive to provide maximum information precision $\alpha_{k_H}^* = 1$ in the baseline, but the semi-pool equilibrium under type uncertainty provides $\alpha_{\theta, k_H}^* = 1$, as a result of advisor's pooling and choice of adjusting from only low type information. This is exactly opposite with the above situation where a good type advisor moves from full to partial truthtelling.

Yet, the conclusion on high type does not necessarily extend to aggregate efficiency. For the first example, ending in case 2 equilibrium results in $k_H + k_L \leq -2 \leq 0$ and so $\frac{1}{4} - \frac{1}{2k_L} < \frac{1}{4} - \frac{1}{k_H + k_L}$ and $\frac{1}{2} < \frac{1}{4} - \frac{1}{k_H + k_L}$, and so low type information is not necessarily hurt by conflict uncertainty. The inequality's left hand side is the low type precision in case 1 equilibrium, and the right hand side is the low type precision in case 2 interior solution under type uncertainty. However, the aggregate efficiency for such case 2 is bounded above at $\frac{1}{4} + \frac{1}{2} = \frac{3}{4}$, which means the aggregate efficiency is definitely lower for conflict uncertainty, only when the low type's information under no uncertainty is greater than $\frac{1}{2}$, i.e. $k_L \geq -2$, the low type's conflict is restricted. Alternatively, we can end up with a lower aggregate efficiency under no uncertainty in the second example if k_H and k_L are too close such that both $k \in (-2, -\frac{2}{3})$.

Then, it stresses again the importance of understanding the absolute and relative magnitude of conflict when assessing whether conflict uncertainty is beneficial for advisor's information provision. Both are crucial to determine whether resolving conflict uncertainty is a good idea, as it pins down which two equilibria the advisor is facing for this cross comparison. Due to time constraints, I only analyse a subset of such comparison. From these comparisons, we already discover an interaction between conflict disparity of both types and the magnitude for low/high type alone matters. When conflict disparity is large, ensuring low type advisor has

limited conflict and resolve conflict uncertainty can be good. When conflict disparity is small, the conclusion would depend on whether resolving uncertainty for the high type indeed realises full information precision. If after resolving conflicts, advisor is in a conflict-revealling equilibrium (signals private benefit) but the high type does not provide $\alpha_{k_H}^* = 1$, then, revealing conflict might instead reduce the signal's aggregate efficiency (measured in ex ante terms). Then, award good type advisors when incentive indeed sufficiently aligns with advisors becomes crucial to realise benefit from resolving advisor's conflict. In addition, encourage activities that has some degree of fully alignment with information sales, regardless of uncertainties in further conflict. Awarding advisors once know that incentives are sufficiently aligned (increasing value of k_H) without type uncertainty would not improve aggregate efficiency, but might help when advisor conflict is uncertain when charging the fee. These results provide more basis to the mixed effects of conflict resolving policies ex ante, (Li and Madarász (2008), Frankel and Kartik (2019), Gesche (2021)). Also, they correspond to some of the results emerged from these models in cheap-talk context, such as agents can generate some more degree of communication even under type uncertainty. However, such results are not easily obtained in models where information is sold to clients.

Chapter 6

Conclusion

This paper presents a financial advisory relationship in which the advisor has a conflict of interest related to the client's ex post information quality, while such signal performance is not ex ante verifiable or contractible. Even though conventional wisdom suggests that limiting an advisory firm's conflicts, raising agents' reputation costs for providing low-quality information, or exposing conflicts beforehand promotes information provision, the empirical, theoretical, and experimental literature sometimes does not agree with each other. This paper offers an explanation for these mixed results.

When the advisor is determining trade-offs between providing more information to the client and enhancing private benefit by manipulating information quality, such trade-offs depend on whether the advisor understands the exact degree of conflict at the time the advisory decision is made. If there is no uncertainty about the conflict scale, the advisor provides what she knows with sufficient incentive alignment, and information quality is monotonically decreasing with the conflict scale. The advisor signalling the conflict does not imply full information, as exposing the conflict only indicates a relative, not absolute, value of the private-benefit scale. In this case, mandatory disclosure can be helpful only if paired with restricting the advisor's conflict of interest.

However, when the advisor has uncertainty about her conflict, the trade-off faced becomes between information quality and ex ante private benefit. She faces multiple equilibria in signal structures, which can lead to non-monotonic informa-

tion quality contingent on the conflict scale. Restricting conflicts of interest might backfire on information quality when the advisor's equilibrium decision changes from one potential equilibrium to another. As those equilibria follow different decision rules compared with the case in which the advisor knows the conflict, it also becomes hard to conclude that resolving such conflict uncertainty necessarily promotes information quality.

To assess the impacts of restricting conflicts of interest, mandatory disclosure, and the extent to which advisors voluntarily disclose their conflicts through fees charged, one needs to measure the degree of conflict across the advisory market in order to pin down comparisons and arrive at a clear result. This work also calls for careful consideration before introducing such regulations. Other paths addressing limitation of this work, including introducing competition among advisors, allowing clients to have other source of private information, would be interesting to work on.

Appendix A

Solve the price-setting process

The market maker sets price according to $p_d = \mathbb{E}(\theta \mid d, R_{\theta,k})$. As $\theta \in \{0, 1\}$, $p_d = \mathbb{P}(\theta = 1 \mid d, R_{\theta,k})$. Under the assumption that if not bought the signal, $I = 0$,¹ we have the following:

- At $d = \{-2, 2\}$, client bought the signal and buys/sells.
- At $d = \{-1, 1\}$, client did not buy the signal and the aggregate order flow comes from the noise trader. Then, MM prices asset at prior.
- At $d = 0$, MM is confounded by the order flow and does not know informed's identity. This means MM could face $\{I = 1, U = -1\}$ or $\{I = -1, U = 1\}$, due to uninformed can only buy or sell.

Then, the problem is to work out p_d and the client's strategies (trading unit and probability with which the client follows the signal Y). To simplify notations, it is easier to describe the client's strategies in their inferred advisor's type $\hat{k}$, which is determined in time 1: $\left\{I, \mu_{Y,\hat{k}}, \mathbb{P}(I = 0 \mid Y, \hat{k})\right\}$, and the value for expected trading

¹Justification as in the main text: the client knows she has no information and is a strategic agent.

profit π .

$$\begin{aligned}
p_2 &= \mathbb{P}(\theta = 1 | d = 2, R_{\theta,k}) \\
&= \mathbb{P}(\theta = 1 | I = 1, U = 1, R_{\theta,k}) \\
&\quad (\text{independent uninformed action } U) \\
&= \mathbb{P}(\theta = 1 | I = 1, R_{\theta,k}) \\
&= \frac{\mathbb{P}(\theta = 1 | R_{\theta,k}) \mathbb{P}(I = 1 | \theta = 1, R_{\theta,k})}{\mathbb{P}(\theta = 1 | R_{\theta,k}) \mathbb{P}(I = 1 | \theta = 1, R_{\theta,k}) + \mathbb{P}(\theta = 0 | R_{\theta,k}) \mathbb{P}(I = 1 | \theta = 0, R_{\theta,k})}
\end{aligned}$$

Similarly,

$$\begin{aligned}
p_{-2} &= \mathbb{P}(\theta = 1 | d = -2, R_{\theta,k}) \\
&= \frac{\mathbb{P}(\theta = 1 | R_{\theta,k}) \mathbb{P}(I = -1 | \theta = 1, R_{\theta,k})}{\mathbb{P}(\theta = 1 | R_{\theta,k}) \mathbb{P}(I = -1 | \theta = 1, R_{\theta,k}) + \mathbb{P}(\theta = 0 | R_{\theta,k}) \mathbb{P}(I = -1 | \theta = 0, R_{\theta,k})}
\end{aligned}$$

And at $d = 0$,

$$\begin{aligned}
p_0 &= \mathbb{P}(\theta = 1 | d = 0, R_{\theta,k}) \\
&= \frac{\mathbb{P}(\theta = 1 | R_{\theta,k}) \mathbb{P}(d = 0 | \theta = 1, R_{\theta,k})}{\mathbb{P}(\theta = 1 | R_{\theta,k}) \mathbb{P}(d = 0 | \theta = 1, R_{\theta,k}) + \mathbb{P}(\theta = 0 | R_{\theta,k}) \mathbb{P}(d = 0 | \theta = 0, R_{\theta,k})}
\end{aligned}$$

where

$$\begin{aligned}
\mathbb{P}(d = 0 | \theta = 1, R_{\theta,k}) &= \mathbb{P}(I = 1, U = -1 | \theta = 1, R_{\theta,k}) \\
&\quad + \mathbb{P}(I = -1, U = 1 | \theta = 1, R_{\theta,k}) \\
&= \mathbb{P}(I = 1 | \theta = 1, R_{\theta,k}) \mathbb{P}(U = -1) \\
&\quad + \mathbb{P}(I = -1 | \theta = 1, R_{\theta,k}) \mathbb{P}(U = 1)
\end{aligned}$$

and similarly

$$\begin{aligned}
\mathbb{P}(d = 0 | \theta = 0, R_{\theta,k}) &= \mathbb{P}(I = 1 | \theta = 0, R_{\theta,k}) \mathbb{P}(U = -1) \\
&\quad + \mathbb{P}(I = -1 | \theta = 0, R_{\theta,k}) \mathbb{P}(U = 1)
\end{aligned}$$

To simplify the price expressions, we need to know the exact value for $\mathbb{P}(\theta|R_{\theta,k})$, which means to figure out what happens in time 1. However, at time 2, we know that without assuming any equilibrium structure of $R_{\theta,k}$, the belief $\mathbb{P}(\theta|R_{\theta,k})$ is the function of $\mathbb{P}(\theta)$ and $\mathbb{P}(k)$ by Bayesian updating, and $\mathbb{P}(\theta|R_{\theta,k}) \in (0, 1)$ otherwise the client can free ride information completely.

To work out fully the client's strategies, first, we need to know that conditional on an inferred advisor's type $\hat{k}$ (given a structure of $R_{\theta,k}$), the respective conditional probabilities of buy and sell. I denote the probability $\mu_{Y,\hat{k}}$ with which client follows signal Y with advisor's inferred type $\hat{k}$. Then, we can write:

$$\begin{aligned}\mathbb{P}(I = 1|\theta = 1, \hat{k}) &= \mathbb{P}(I = 1|Y = 1, \hat{k})\mathbb{P}(Y = 1|\theta = 1, \hat{k}) \\ &\quad + \mathbb{P}(I = 1|Y = 0, \hat{k})\mathbb{P}(Y = 0|\theta = 1, \hat{k}) \\ &= \mu_{1,\hat{k}}\hat{\alpha}_{1,\hat{k}} + (1 - \mu_{0,\hat{k}} - \mathbb{P}(I = 0|\theta = 1, \hat{k}, Y))(1 - \hat{\alpha}_{1,\hat{k}})\end{aligned}$$

$$\begin{aligned}\mathbb{P}(I = 1|\theta = 0, \hat{k}) &= \mathbb{P}(I = 1|Y = 1, \hat{k})\mathbb{P}(Y = 1|\theta = 0, \hat{k}) \\ &\quad + \mathbb{P}(I = 1|Y = 0, \hat{k})\mathbb{P}(Y = 0|\theta = 0, \hat{k}) \\ &= \mu_{1,\hat{k}}(1 - \hat{\alpha}_{0,\hat{k}}) + (1 - \mu_{0,\hat{k}} - \mathbb{P}(I = 0|\theta = 1, \hat{k}, Y))\hat{\alpha}_{0,\hat{k}}\end{aligned}$$

and

$$\begin{aligned}\mathbb{P}(I = -1|\theta = 1, \hat{k}) &= \mathbb{P}(I = -1|Y = 1, \hat{k})\mathbb{P}(Y = 1|\theta = 1, \hat{k}) \\ &\quad + \mathbb{P}(I = -1|Y = 0, \hat{k})\mathbb{P}(Y = 0|\theta = 1, \hat{k}) \\ &= (1 - \mu_{1,\hat{k}} - \mathbb{P}(I = 0|\theta = 1, \hat{k}, Y))\hat{\alpha}_{1,\hat{k}} + \mu_{0,\hat{k}}(1 - \hat{\alpha}_{1,\hat{k}})\end{aligned}$$

$$\begin{aligned}\mathbb{P}(I = -1|\theta = 0, \hat{k}) &= \mathbb{P}(I = -1|Y = 1, \hat{k})\mathbb{P}(Y = 1|\theta = 0, \hat{k}) \\ &\quad + \mathbb{P}(I = -1|Y = 0, \hat{k})\mathbb{P}(Y = 0|\theta = 0, \hat{k}) \\ &= (1 - \mu_{1,\hat{k}} - \mathbb{P}(I = 0|\theta = 1, \hat{k}, Y))(1 - \hat{\alpha}_{0,\hat{k}}) + \mu_{0,\hat{k}}\hat{\alpha}_{0,\hat{k}}\end{aligned}$$

and also, $\mathbb{P}(I = 1|\theta = 1, R_{\theta,k}) = \sum_{\hat{k}} \mathbb{P}(I = 1|\theta = 1, \hat{k})\mathbb{P}(\hat{k}|R_{\theta,k}, \theta = 1)$, where

$\mathbb{P}(\hat{k}|R_{\theta,k}, \theta = 1)$ represent making joint inference from $R_{\theta,k}$ to types θ, k , and without any assumptions on exact equilibrium structures, we does not know whether the inferred type $\hat{k}$ and $\theta = 1$ happens at the same time. Like before, $\mathbb{P}(\hat{k}|R_{\theta,k}, \theta = 1)$ is only related to priors $\mathbb{P}(\theta)$ and $\mathbb{P}(k)$ by Bayesian updating. As $\mathbb{P}(\hat{k}|R_{\theta,k}, \theta = 1) \geq 0$ and the strict inequality holds for at least one of the advisor's inferred type $\hat{k}$, $\mathbb{P}(I = 1|\theta = 1, R_{\theta,k})$ increases in $\mathbb{P}(I = 1|\theta = 1, \hat{k})$.

Then, write out expected trading profit π when signal:

$$\begin{aligned}
\pi &= \mathbb{E}(I(\theta - \hat{p}_d)|R_{\theta,k}) \\
&= \mathbb{E}_{\theta}(\mathbb{E}(I(\theta - \hat{p}_d)|\theta, R_{\theta,k})) \\
&= \sum_{\theta \in \{0,1\}} \mathbb{P}(\theta | R_{\theta,k}) \mathbb{E}(I(\theta - p_d) | \theta, R_{\theta,k}) \\
&= \mathbb{P}(\theta = 1 | R_{\theta,k}) \mathbb{E}[(1 - p_d)(2\mathbb{P}(I = 1|\theta = 1, R_{\theta,k}) - 1)] \\
&\quad + \mathbb{P}(\theta = 0 | R_{\theta,k}) \mathbb{E}[p_d(2\mathbb{P}(I = -1|\theta = 0, R_{\theta,k}) - 1)]
\end{aligned}$$

To maximise π , informed trader maximises the probabilities of correct trades ($\mathbb{P}(I = 1|\theta = 1, R_{\theta,k})$ and $\mathbb{P}(I = -1|\theta = 0, R_{\theta,k})$, bring weakly positive realised profits) and minimise probabilities of incorrect trades ($\mathbb{P}(I = 1|\theta = 0, R_{\theta,k})$ and $\mathbb{P}(I = -1|\theta = 1, R_{\theta,k})$, bring weakly negative profits). When $\alpha_{\theta,k} \geq \frac{1}{2}$, the signal is at least weakly informative. If $\alpha_{\theta,k} > \frac{1}{2}$, $1 - \alpha_{\theta,k} < \frac{1}{2} < \alpha_{\theta,k}$, the client puts maximum weight on the term $\alpha_{\theta,k}$ that is larger for the correct trading probabilities and puts maximum weight on the term $1 - \alpha_{\theta,k}$ for the incorrect trading probabilities. This involves setting $\mu_{Y,\hat{k}} = 1$ and $\mathbb{P}(I = 0|\theta, \hat{k}, Y) = 0$, i.e. equilibrium $I = \{-1, 1\}$. When $\alpha_{\theta,k} = \frac{1}{2}$, apply a tie-breaking rule that client in $t = 2$ follows signal, i.e. $\mu_{Y,\hat{k}} = 1$, whenever indifferent between choices of $\mu_{Y,\hat{k}} \in [0, 1]$. The client follows signal in probability 1 translate to the main texts $\mathbb{P}(I = 1 | Y = 1, R_{\theta,k}) = 1$.

Then, the prices can be simplified as:

$$\begin{aligned}
p_2 &= \frac{\sum_{\hat{k}} \hat{\alpha}_{1,\hat{k}} \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 1) \mathbb{P}(\theta = 1 | R_{\theta,k})}{\sum_{\hat{k}} \hat{\alpha}_{1,\hat{k}} \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 1) \mathbb{P}(\theta = 1 | R_{\theta,k}) + \sum_{\hat{k}} (1 - \hat{\alpha}_{0,\hat{k}}) \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 0) \mathbb{P}(\theta = 0 | R_{\theta,k})} \\
p_{-2} &= \frac{\sum_{\hat{k}} (1 - \hat{\alpha}_{1,\hat{k}}) \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 1) \mathbb{P}(\theta = 1 | R_{\theta,k})}{\sum_{\hat{k}} (1 - \hat{\alpha}_{1,\hat{k}}) \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 1) \mathbb{P}(\theta = 1 | R_{\theta,k}) + \sum_{\hat{k}} \hat{\alpha}_{0,\hat{k}} \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 0) \mathbb{P}(\theta = 0 | R_{\theta,k})} \\
p_0 &= \frac{\mathbb{P}(d = 0 | \theta = 1, R_{\theta,k}) \mathbb{P}(\theta = 1 | R_{\theta,k})}{\mathbb{P}(d = 0 | \theta = 1, R_{\theta,k}) \mathbb{P}(\theta = 1 | R_{\theta,k}) + \mathbb{P}(d = 0 | \theta = 0, R_{\theta,k}) \mathbb{P}(\theta = 0 | R_{\theta,k})}
\end{aligned}$$

where

$$\begin{aligned}
\mathbb{P}(d = 0 | \theta = 1, R_{\theta,k}) &= \left(\sum_{\hat{k}} \hat{\alpha}_{1,\hat{k}} \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 1) \right) \mathbb{P}(U = -1) \\
&\quad + \left(\sum_{\hat{k}} (1 - \hat{\alpha}_{1,\hat{k}}) \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 1) \right) \mathbb{P}(U = 1) \\
\mathbb{P}(d = 0 | \theta = 0, R_{\theta,k}) &= \left(\sum_{\hat{k}} (1 - \hat{\alpha}_{0,\hat{k}}) \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 0) \right) \mathbb{P}(U = -1) \\
&\quad + \left(\sum_{\hat{k}} \hat{\alpha}_{0,\hat{k}} \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 0) \right) \mathbb{P}(U = 1)
\end{aligned}$$

From structure of $R_{\theta,k}$, we know $\mathbb{P}(k | R_{\theta,k}, \theta) \in \{0, 1, \frac{1}{2}\}$ and $\mathbb{P}(\theta | R_{\theta,k}) = \frac{1}{2}$. We also know that the corresponding structure of $\alpha_{\theta,k}$ and can compute prices. And then p_0 becomes

$$\begin{aligned}
&\frac{\sum_{\hat{k}} \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 1) \left(\hat{\alpha}_{1,\hat{k}}(1 - \varepsilon) + (1 - \hat{\alpha}_{1,\hat{k}})\varepsilon \right)}{\sum_{\hat{k}} \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 1) \left(\hat{\alpha}_{1,\hat{k}}(1 - \varepsilon) + (1 - \hat{\alpha}_{1,\hat{k}})\varepsilon \right) + \sum_{\hat{k}} \mathbb{P}(\hat{k} | R_{\theta,k}, \theta = 0) \left((1 - \hat{\alpha}_{0,\hat{k}})(1 - \varepsilon) + \hat{\alpha}_{0,\hat{k}}\varepsilon \right)} \\
&= \begin{cases} \hat{\alpha}_{1,\hat{k}}(1 - \varepsilon) + \varepsilon(1 - \hat{\alpha}_{1,\hat{k}}) & \text{if } \mathbb{P}(\hat{k} | R_{\theta,k}, \theta) = 1 \\ \frac{\hat{\alpha}_{1,\hat{k}} + \hat{\alpha}_{1,\hat{k}'}}{2}(1 - \varepsilon) + \varepsilon(1 - \frac{\hat{\alpha}_{1,\hat{k}} + \hat{\alpha}_{1,\hat{k}'}}{2}) & \text{if } \mathbb{P}(\hat{k} | R_{\theta,k}, \theta) = \frac{1}{2} \\ \hat{\alpha}_{1,\hat{k}'}(1 - \varepsilon) + \varepsilon(1 - \hat{\alpha}_{1,\hat{k}'}) & \text{if } \mathbb{P}(\hat{k} | R_{\theta,k}, \theta) = 0 \text{ and } \hat{k}' \neq \hat{k} \end{cases}
\end{aligned}$$

Similarly,

$$p_2 = \begin{cases} \hat{\alpha}_{1,\hat{k}} & \text{if } \mathbb{P}(\hat{k}|R_{\theta,k}, \theta) = 1 \\ \frac{\hat{\alpha}_{1,\hat{k}} + \hat{\alpha}_{1,\hat{k}'}}{2} & \text{if } \mathbb{P}(\hat{k}|R_{\theta,k}, \theta) = \frac{1}{2} \\ \hat{\alpha}_{1,\hat{k}'} & \text{if } \mathbb{P}(\hat{k}|R_{\theta,k}, \theta) = 0 \text{ and } \hat{k}' \neq \hat{k} \end{cases}$$

and

$$p_{-2} = \begin{cases} 1 - \hat{\alpha}_{1,\hat{k}} & \text{if } \mathbb{P}(\hat{k}|R_{\theta,k}, \theta) = 1 \\ 1 - \frac{\hat{\alpha}_{1,\hat{k}} + \hat{\alpha}_{1,\hat{k}'}}{2} & \text{if } \mathbb{P}(\hat{k}|R_{\theta,k}, \theta) = \frac{1}{2} \\ 1 - \hat{\alpha}_{1,\hat{k}'} & \text{if } \mathbb{P}(\hat{k}|R_{\theta,k}, \theta) = 0 \text{ and } \hat{k}' \neq \hat{k} \end{cases}$$

and with each belief structure determined in time 1 action, plugged in prices and beliefs to get the expected profit π .

Appendix B

Proof for Lemma 1

Proof of Lemma 1. From advisor's problem and the client's IR constraint, the problem can be simplified into optimisation according to a single decision variable $\alpha_{\theta,k}$. This is summarised by below: Client's IR constraint must be binding in equilibrium, i.e. equilibrium expected trading profit π and fee $R_{\theta,k}$ are equal. The proof is as follows. The advisor is never optimal to charge below expected trading profit by a constant, say m , as improving $R_{\theta,k} = \pi - m$ by a smaller amount $m' < m$ always dominates. Then, suppose that advisor relinquish a fraction of trading profit. Note that conditional on signal structure, profit should be function of at most two distinct precision variables $\alpha_{\theta,k}$ and $\alpha_{\theta,k'}$. Expected trading profit π is expectation of trading gains using precision $\{\alpha_{\theta,k}, \alpha_{\theta,k'}\}$ with weight the belief attached to state $\mathbb{P}(k | R_{\theta,k}, \theta)$, and name such trading profit $\pi_{\theta,k}$. Suppose that $\mathbb{P}(k | R_{\theta,k}, \theta) \in \{0, 1\}$ and advisor charges a share of profit $\gamma_{\theta,k}$ on the expected trading profit conditional on signal structure with precision $\alpha_{\theta,k}$. When $\alpha_{\theta,k} \geq \frac{1}{2}$, profit $\pi_{\theta,k}$ increases in $\gamma_{\theta,k}$, and so $\gamma_{\theta,k} = 1$. Any deviation to $\gamma_{\theta,k} < 1$ would be dominated by $\tilde{\gamma}_{\theta,k}$, an infinitesimal amount larger than $\gamma_{\theta,k} < 1$.¹ The remaining case is $\mathbb{P}(k | R_{\theta,k}, \theta) \in (0, 1)$, and the same argument holds, as this probability is not a function of $\alpha_{\theta,k}$ but only the prior probabilities $\mathbb{P}(\theta)$, $\mathbb{P}(k)$. Applying the reasoning to both $\alpha_{\theta,k}$ and $\alpha_{\theta,k'}$, advisor still want to charge $R_{\theta,k} = \pi$.

¹Otherwise, advisor wants to charge $\gamma_{\theta,k} = 0$. However, this implies client has bought strictly uninformative signals and expected trading profit decreases, not reflected in the advisor's information price. Then, client's IR is violated as a strict loss would have been made if bought the signal. That also implies advisor cannot credibly send $\alpha_{\theta,k} < \frac{1}{2}$.

Appendix C

Appendices for analysis sections

Checking equilibria in section 5, semi-pool equilibria. The process start by marking

$$\frac{\partial EU_A}{\partial \alpha_{\theta,k}} = \frac{1}{2} \varepsilon (1 - \varepsilon) (2 + k' \alpha_{\theta,k'} + 2k \alpha_{\theta,k} + k(\alpha_{\theta,k'} - 1))$$

The solution $(\alpha_{\theta,k}^* \in (\frac{1}{2}, 1), \alpha_{\theta,k'} = 1)$ requires

$$\begin{aligned} \frac{1}{2} \varepsilon (1 - \varepsilon) (2 + k' + 2k \alpha_{\theta,k}^*) &= 0 \\ \frac{1}{2} \varepsilon (1 - \varepsilon) (2 + k' \alpha_{\theta,k'} + 2k \alpha_{\theta,k}^*) &> 0 \text{ for } \alpha_{\theta,k'} \in \left[\frac{1}{2}, 1\right] \\ \varepsilon (1 - \varepsilon) \alpha_{\theta,k}^* \left(1 + \frac{1}{2}(k \alpha_{\theta,k}^* + k')\right) &\geq 0 \\ \alpha_{\theta,k}^* &\in \left(\frac{1}{2}, 1\right) \end{aligned}$$

Note that $k_L < k_H$, and I only look at cases where $k_L < 0$; otherwise the equilibrium end up with $\alpha_{\theta,k} = 1$ due to incentive alignment. The second-order condition holds for $k < 0$, where we obtain an interior solution. From the first and last condition, it becomes clear that we put $\alpha_{\theta,k_L}^* = -\frac{2+k_H}{2k_L}$, $\alpha_{\theta,k_H}^* = 1$, $k = k_L$. The other two conditions check whether those solutions gives nonnegative utilities and precision in range. That results in

$$k_H \geq -2 \text{ and } k_L < 0 \text{ and } k_H + 2 + k_L > 0 \text{ and } k_H + 2 + 2k_L < 0$$

The solution $\left(\alpha_{\theta,k}^* = \frac{1}{2}, \alpha_{\theta,k'} = 1\right)$ requires

$$\begin{aligned} \frac{1}{2}\varepsilon(1-\varepsilon)(2+k'\alpha_{\theta,k'}+k) &> 0 \text{ for } \alpha_{\theta,k'} \in \left[\frac{1}{2}, 1\right] \\ \frac{1}{2}\varepsilon(1-\varepsilon)(2+k'+2k\alpha_{\theta,k}) &< 0 \text{ for } \alpha_{\theta,k} \in \left[\frac{1}{2}, 1\right] \\ \varepsilon(1-\varepsilon)\frac{1}{2}\left(1+\frac{1}{2}\left(\frac{1}{2}k+k'\right)\right) &\geq 0 \end{aligned}$$

It becomes clear that we put $\alpha_{\theta,k_L}^* = \frac{1}{2}$, $\alpha_{\theta,k_H}^* = 1$, $k = k_L$. The last condition check individual rationality constraint, which result in

$$2 + \frac{1}{2}k_L + k_H \geq 0$$

The first two conditions check whether those solutions are incentive compatible. As those are linear constraints, we only need to check at boundaries, i.e. at precision $\frac{1}{2}$ or 1. These result in

$$k_L < 0 \text{ and } k_H + 2 + k_L < 0 \text{ and } k_H + 2 + \frac{1}{2}k_L > 0$$

The other two solutions do not work for the following reasons. The pair of precision $\left(\alpha_{\theta,k} = \frac{1}{2}, \alpha_{\theta,k'} \in \left(\frac{1}{2}, 1\right)\right)$ requires $\alpha_{\theta,k_H} = \frac{-4+k_L-k_H}{4k_L}$. For the pairs of (k_L, k_H) that $\alpha_{\theta,k_H}^* \in \left(\frac{1}{2}, 1\right)$, the partial in α_{θ,k_L} is not guaranteed negative. The two interior solutions $\alpha_{\theta,k}^* = \frac{2+k'}{k'-k}$ which does not satisfy precision bound constraints $\alpha_{\theta,k}^* \in \left(\frac{1}{2}, 1\right)$ for joint (k, k') . Intuitively those combination where $\alpha_H < 1$ does not fit advisor's incentive compatibility constraints.

Supplementary graph for section 5.3.3. See fig. C.1 for a graphical representation for regions (in red) when advisor's utility is higher in the fully pool equilibrium in section 5.1.

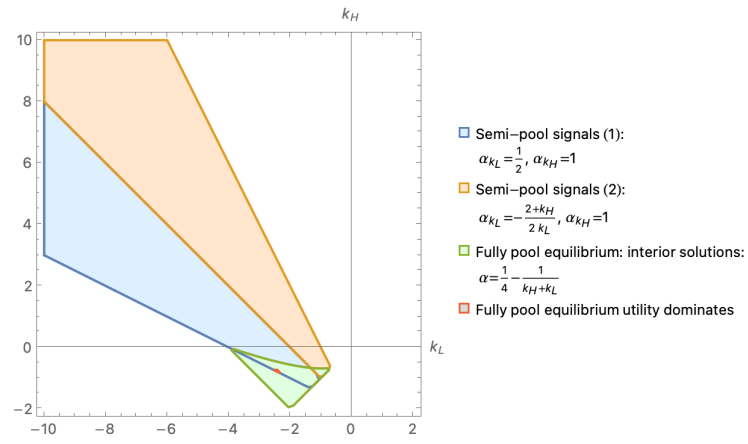


Figure C.1: Region When Fully Pool Equilibrium Achieves a Higher Advisor Utility

Bibliography

- A. Admati and P. Pfleiderer. Selling and trading on information in financial markets. *American Economic Review*, 78(2):96–103, 1988.
- A. R. Admati and P. Pfleiderer. A monopolistic market for information. *Journal of Economic Theory*, 39(2):400–438, 1986.
- A. R. Admati and P. Pfleiderer. Does it all add up? benchmarks and the compensation of active portfolio managers. *The Journal of Business*, 70(3):323–350, 1997.
- G.-M. Angeletos, G. Lorenzoni, and A. Pavan. Wall Street and Silicon Valley: A Delicate Interaction. *The Review of Economic Studies*, 90(3):1041–1083, 07 2022.
- S. Banerjee and B. Breon-Drish. Strategic trading and unobservable information acquisition. *Journal of Financial Economics*, 138(2):458–482, 2020.
- S. Banerjee and B. Green. Signal or noise? uncertainty and learning about whether other traders are informed. *Journal of Financial Economics*, 117(2):398–423, 2015.
- J. Bao, M. O’Hara, and X. (Alex) Zhou. The Volcker Rule and corporate bond market making in times of stress. *Journal of Financial Economics*, 130(1):95–113, 2018.
- S. Behnk, I. Barreda-Tarrazona, and A. García-Gallego. The role of ex post transparency in information transmission—an experiment. *Journal of Eco-*

- nomic Behavior & Organization*, 101:45–64, 2014. ISSN 0167-2681. doi: <https://doi.org/10.1016/j.jebo.2014.02.006>.
- R. Benabou and G. Laroque. Using privileged information to manipulate markets: Insiders, gurus, and credibility. *The Quarterly Journal of Economics*, 107(3): 921–958, 1992.
- U. Bhattacharya, A. Hackethal, S. Kaesler, B. Loos, and S. Meyer. Is Unbiased Financial Advice to Retail Investors Sufficient? Answers from a Large Field Study. *The Review of Financial Studies*, 25(4):975–1032, 2012.
- P. Bolton, X. Freixas, and J. Shapiro. Conflicts of interest, information provision, and competition in the financial services industry. *Journal of Financial Economics*, 85(2):297–330, 2007.
- P. Bolton, X. Freixas, and J. Shapiro. The credit ratings game. *The Journal of Finance*, 67(1):85–111, 2012.
- R. Calcagno and C. Monticone. Financial literacy and the demand for financial advice. *Journal of Banking & Finance*, 50:363–380, 2015.
- G. Cespa. Information sales and insider trading with long-lived information. *The Journal of Finance*, 63(2):639–672, 2008.
- A. Chakraborty and B. Yılmaz. Manipulation in market order models. *Journal of Financial Markets*, 7(2):187–206, 2004.
- J. Chalmers and J. Reuter. Is conflicted investment advice better than no advice? *Journal of Financial Economics*, 138(2):366–387, 2020.
- B. Chang and M. Szydlowski. The market for conflicted advice. *Journal of Finance*, 75(2):867–903, 2020.
- D. Cuoco and R. Kaniel. Equilibrium prices in the presence of delegated portfolio management. *Journal of Financial Economics*, 101(2):264–296, 2011. ISSN 0304-405X.

- A.-I. De Moragas. Disclosing decision makers' private interests. *European Economic Review*, 150:104282, 2022.
- J. Dow and G. Gorton. Stock market efficiency and economic efficiency: Is there a connection? *The Journal of Finance*, 52(3):1087–1129, 1997.
- J. Dow, I. Goldstein, and A. Guembel. Incentives for Information Production in Markets where Prices Affect Real Investment. *Journal of the European Economic Association*, 15(4):877–909, 02 2017.
- D. Duffie. Market making under the proposed volcker rule. Unpublished Working Paper, 2012.
- M. Egan. Brokers versus retail investors: Conflicting interests and dominated products. *The Journal of Finance*, 74(3):1217–1260, 2019.
- A. Frankel and N. Kartik. Muddled information. *Journal of Political Economy*, 127(4):1739–1776, 2019.
- D. García and F. Sangiorgi. Information sales and strategic trading. *Review of Financial Studies*, 24(9):3069–3104, 2011.
- T. Gesche. De-biasing strategic communication. *Games and Economic Behavior*, 130:452–464, 2021.
- I. Goldstein and A. Guembel. Manipulation and the allocational role of prices. *The Review of Economic Studies*, 75(1):133–164, 2008.
- I. Goldstein, E. Ozdenoren, and K. Yuan. Trading frenzies and their impact on real investment. *Journal of Financial Economics*, 109(2):566–582, 2013. ISSN 0304-405X.
- J. Golec and L. Starks. Performance fee contract change and mutual fund risk. *Journal of Financial Economics*, 73(1):93–118, 2004.
- A. Hackethal, M. Haliassos, and T. Jappelli. Financial advisors: A case of babysitters? *Journal of Banking & Finance*, 36(2):509–524, 2012.

- D. Hoechle, S. Ruenzi, N. Schaub, and M. Schmid. Financial Advice and Bank Profits. *The Review of Financial Studies*, 31(11):4447–4492, 2018. ISSN 0893-9454.
- R. Inderst and M. Ottaviani. Misselling through agents. *American Economic Review*, 99(3):883–908, 2009.
- R. Inderst and M. Ottaviani. Competition through commissions and kickbacks. *American Economic Review*, 102(2):780–809, 2012a.
- R. Inderst and M. Ottaviani. Financial advice. *Journal of Economic Literature*, 50(2):494–512, 2012b.
- H. Ismayilov and J. Potters. Disclosing advisor’s interests neither hurts nor helps. *Journal of Economic Behavior & Organization*, 93:314–320, 2013. ISSN 0167-2681. doi: <https://doi.org/10.1016/j.jebo.2013.03.034>. URL <https://www.sciencedirect.com/science/article/pii/S0167268113000796>.
- R. A. Jarrow. Market manipulation, bubbles, corners, and short squeezes. *The Journal of Financial and Quantitative Analysis*, 27(3):311–336, 1992.
- M. Kartal and J. Tremewan. An offer you can refuse: The effect of transparency with endogenous conflict of interest. *Journal of Public Economics*, 161:44–55, 2018. ISSN 0047-2727.
- C. W. Li and A. Tiwari. Incentive Contracts in Delegated Portfolio Management. *The Review of Financial Studies*, 22(11):4681–4714, 2009.
- M. Li and K. Madarász. When mandatory disclosure hurts: Expert advice and conflicting interests. *Journal of Economic Theory*, 139(1):47–74, 2008.
- X. Li and W. Wu. Portfolio pumping and fund performance ranking: A performance-based compensation contract perspective. *Journal of Banking & Finance*, 105:94–106, 2019.

- L. Ma, Y. Tand, and J.-P. Gómez. Portfolio manager compensation in the u.s. mutual fund industry. *The Journal of Finance*, 74(2):587–638, 2019.
- A. Malenko and N. Malenko. Proxy advisory firms: The economics of selling information to voters. *Journal of Finance*, 74(5):2441–2490, 2019.
- J. Morgan and P. C. Stocken. An analysis of stock recommendations. *RAND Journal of Economics*, 34(1):183–203, 2003.
- M. Ottaviani and P. N. Sørensen. Professional advice. *Journal of Economic Theory*, 126(1):120–142, 2006.
- S. A. Ross. Compensation, incentives, and the duality of risk aversion and riskiness. *The Journal of Finance*, 59(1):207–225, 2004.
- J. Sotes-Paladino and F. Zapatero. Carrot and stick: A role for benchmark-adjusted compensation in active fund management. *Journal of Financial Intermediation*, 52:100981, 2022.
- L. T. Starks. Performance incentive fees: An agency theoretic approach. *The Journal of Financial and Quantitative Analysis*, 22(1):17–32, 1987.
- N. M. Stoughton, Y. Wu, and J. Zechner. Intermediated investment management. *The Journal of Finance*, 66(3):947–980, 2011.
- L. Veldkamp. Information markets and the comovement of asset prices. *Review of Economic Studies*, 73(3):823–845, 2006.
- Y. Xiong and L. Yang. Secret and Overt Information Acquisition in Financial Markets. *The Review of Financial Studies*, 36(9):3643–3692, 2023.