

Discussion Paper Series

IZA DP No. 18832

July 2026

Identifying Level- k Reasoning in Repeated Games: Strategies, Beliefs, and Cognitive Ability

David Gill

Purdue University
and IZA@LISER

Yaroslav Rosokha

Purdue University

The IZA Discussion Paper Series (ISSN: 2365-9793) ("Series") is the primary platform for disseminating research produced within the framework of the IZA@LISER Network, an unincorporated international network of labour economists coordinated by the Luxembourg Institute of Socio-Economic Research (LISER). The Series is operated by LISER, a Luxembourg public establishment (établissement public) registered with the Luxembourg Business Registers under number J57, with its registered office at 11, Porte des Sciences, 4366 Esch-sur-Alzette, Grand Duchy of Luxembourg.

Any opinions expressed in this Series are solely those of the author(s). LISER accepts no responsibility or liability for the content of the contributions published herein. LISER adheres to the European Code of Conduct for Research Integrity. Contributions published in this Series present preliminary work intended to foster academic debate. They may be revised, are not definitive, and should be cited accordingly. Copyright remains with the author(s) unless otherwise indicated.

Identifying Level- k Reasoning in Repeated Games: Strategies, Beliefs, and Cognitive Ability*

Abstract

In this paper, we identify level- k reasoning in repeated games that operates at the level of a supergame strategy. First, we develop a model of level- k reasoning that incorporates choices over strategies as well as beliefs about the strategies chosen by others. Then, jointly using elicited strategies and beliefs about strategies from the indefinitely repeated Prisoner's Dilemma, and applying a new notion of balance-optimal error tolerance, we classify a substantial fraction of subjects as level-1 or level-2. Moreover, we show that when level- k reasoning operates at the level of a repeated-game strategy, cognitive ability and experience predict higher level reasoning.

JEL classification

C73, D83, D91

Keywords

level- k , repeated game, Prisoner's Dilemma, strategy, beliefs, cognitive ability, experience, elicitation, bounded rationality, experiment, game theory

Corresponding author

David Gill

dgill@econ@gmail.com

* We thank: Shabnam Rahimfallah, Rui Tan and Zhongheng Qiao for excellent research assistance; seminar participants and colleagues for valuable comments and discussion; the Vernon Smith Experimental Economics Laboratory at Purdue University for hosting our experiment; and the Blake Family Fund for Ethics, Leadership and Governance (administered by the Daniels School of Business) for providing funding. This paper supersedes Section 3.3 of Gill and Rosokha (2020), which contained a preliminary version of some of the analysis here, and which was not included or referenced by Gill and Rosokha (2024).

I Introduction

The level- k model of boundedly rational reasoning is a powerful tool for analyzing strategic interactions. The model eschews equilibrium thinking; instead, level- k reasoning anchors beliefs in an instinctive reaction to the game and then adjusts them via iterated best responses (Crawford et al., 2013). Specifically, the level-0 type is strategically unsophisticated, the level-1 type best responds to the belief that all others are level-0, the level-2 type best responds to the belief that all others are level-1, and so on (Nagel, 1995). Alaoui and Penta (2016, 2022)’s model of endogenous depth of reasoning microfound choice of level, while Alós-Ferrer and Buckenmaier (2021) and Gill and Prowse (2023) find that higher level- k types think for longer. The level- k model has been used to help understand strategic behavior in a variety of settings, including auctions (e.g., Crawford and Iriberri, 2007a), undercutting games (e.g., Arad and Rubinstein, 2012), Centipede games (e.g., García-Pola et al., 2020), Cournot games (e.g., Ho et al., 2021), and gas station pricing (Han et al., 2026).

Despite the central role that beliefs play in level- k theory, empirical applications have not used choices and elicited beliefs together to estimate models of level- k behavior.¹ In this paper, we use choice and belief data jointly to identify level- k reasoning in the indefinitely repeated Prisoner’s Dilemma (IRPD), which captures well the trade-off between the short-term payoff from exploiting partners and the long-term gain from building enduring cooperative relationships. Strategies in the IRPD incorporate various forms of sophisticated behavior, including conditional cooperation, punishment, and forgiveness. Given this complexity, we show that data on beliefs are critical for level- k identification because strategy choices alone cannot clearly distinguish different levels, or indeed level- k reasoning from other models of behavior.

We measure level- k reasoning that operates at the level of a supergame strategy (i.e., a plan of action for every round of the repeated game that can condition on the history of play). By contrast, the existing literature that measures level- k reasoning in repeated

¹There are two main reasons for this. First, the existing literature often studies relatively simple games in which different levels generate distinct choice predictions, making it possible to infer levels from choices alone. Second, accommodating noise in both beliefs and choices, and then jointly mapping these two sources of data into levels of reasoning, is methodologically challenging.

games models level- k behavior at the level of a single round or repetition.² We follow a recent stream of papers that directly elicits strategies in the IRPD (Romero and Rosokha, 2018, 2023, Cason and Mui, 2019, Dal Bó and Fréchette, 2019). Specifically, we use the IRPD dataset from our earlier study, Gill and Rosokha (2024).³ In that experiment we elicit supergame strategies and beliefs about strategies: subjects choose a strategy at the beginning of each supergame, rather than playing via direct response (with a sequence of round-by-round decisions), and report beliefs over the strategies chosen by others (in the first and final supergames).

Our dataset is particularly well suited to identifying level- k reasoning that operates at the level of a strategy because it contains both strategy choices and beliefs. Access to data on beliefs is essential for identification, for two reasons. First, different level- k types may choose the same strategy while holding different beliefs. For example, two subjects may choose Always Defect (AD), yet one may do so because she believes opponents randomize uniformly across strategies (holding level-1 beliefs), while the other may do so because she believes opponents will choose AD (holding level-2 beliefs). In this case, observing beliefs allows us to distinguish between reasoning types despite identical strategy choices. Second, the same strategy choice could arise from level- k reasoning or from some other process (e.g., confusion, imitation, efficiency seeking, equilibrium thinking). In this case, observing beliefs allows us to assess whether a subject is consistent with the level- k model. That is, beliefs help identify the reasoning behind chosen strategies, thereby strengthening identification in a belief-driven model like level- k . In Section V, we develop these points further, using specific experimental subjects as examples.

²Some papers estimate level- k reasoning separately in each round (Nagel, 1995, Duffy and Nagel, 1997, Danz et al., 2012, Lindner and Sutter, 2013, Benndorf et al., 2017, Ferraz et al., 2025). Others use data from multiple rounds to estimate level- k reasoning (Stahl, 1996, Ho et al., 1998, Gill and Prowse, 2016, Ho et al., 2021, Castagnetti et al., 2023, Gill and Prowse, 2023, Gill et al., 2025), but the level- k reasoning remains defined at the level of stage-game actions (even when beliefs adjust dynamically, e.g., about the behavior of level-0 types as in Gill and Prowse, 2016). Given this focus, the literature has not measured level- k reasoning in settings such as the IRPD, where round-by-round level- k analysis is silent about depth of reasoning because all types with $k \geq 1$ defect in the stage game.

³Although we use the dataset from Gill and Rosokha (2024), we pursue different empirical goals. Gill and Rosokha (2024) do not structurally estimate types; instead, using only reduced-form analysis, the paper focuses on the relationship between cooperation and beliefs about cooperation, and how both evolve with experience. Gill and Rosokha (2024) find that being matched with a cooperative opponent increases subsequent cooperation and makes beliefs more optimistic. In this paper, we do not study individual-level responses to experience; instead, we use strategies and beliefs about strategies to identify level- k reasoning, and we study how level- k reasoning changes with experience at the aggregate level.

We use elicited strategies and beliefs about strategies instead of direct-response within-supergame round-by-round choices and beliefs. Strategy elicitation is a well-established experimental design that goes back to Selten (1967) and comes with benefits and costs. In our case, the main benefit is that we directly observe strategies and beliefs about strategies. If we used round-by-round data, we would need to structurally estimate the distributions of strategies and beliefs (Aoyagi et al., 2024), and then use posterior probabilities to noisily allocate specific strategies and beliefs to each subject. Furthermore, we study how experience affects level- k reasoning: structural estimation requires the assumption of stable strategy choices and beliefs across many supergames, and so round-by-round data cannot be used to measure level- k reasoning that operates at the level of a strategy before subjects gain experience. Another benefit is that our subjects evaluate strategies only at the beginning of the supergame, as assumed in the theoretical analysis of repeated games.

The main cost is that strategy elicitation could lead to behavior that differs from direct-response play. However, previous evidence and features of our setting help to allay this concern. First, existing evidence shows that when strategies are elicited in the IRPD, the empirical distribution of strategies is similar to that estimated from direct-response play (Romero and Rosokha, 2018, Dal Bó and Fréchette, 2019).⁴ This is consistent with survey evidence in Brandts and Charness (2011) from more general settings (they find no case where the strategy method generates new treatment effects, and that in a majority of surveyed studies behavior is similar under the strategy method and direct response). Second, we emphasize that strategy elicitation captures the dynamic reasoning at the core of repeated-game play. The object of choice remains a repeated-game strategy (a history-contingent plan), with strategies varying in whether and how quickly they punish and relent (see Figure 2). Thus, the decision problem involves reasoning over the full dynamic paths of play implied by pairs of strategies. As in direct-response play, subjects need to consider whether others are likely to start by cooperating or defecting, how harshly defections are likely to be punished, and whether attempting to return to cooperation after a period of punishment is worthwhile. Third, our design keeps the round-by-round representation of the repeated game salient throughout. Before choosing

⁴To enable elicitation of both strategies and beliefs, we restrict attention to 10 strategies. While such a restriction could in principle distort behavior, these strategies include the 8 most common identified via structural estimation using round-by-round choices in the IRPD (Gill and Rosokha, 2024, Section I.D). Consistent with this, our data from elicited strategies replicate key features of direct-response play, although cooperation is broadly stable over supergames when the return to cooperation is high (see Gill and Rosokha, 2024, Appendix VII, for details).

strategies, subjects gained experience with direct-response play in practice supergames played against themselves, and once strategies were chosen, every supergame was played out round-by-round on the subject’s screen similarly to how direct-response games unfold. Fourth, given our experimental parameters, full cooperation (via Grim) and full defection (via Always Defect) are both Nash equilibria of our game with strategy elicitation and subgame-perfect Nash equilibria (SPE) of the direct-response IRPD. Thus we preserve equilibrium multiplicity.⁵

Our first contribution is to develop a model of level- k reasoning in repeated games that operates at the level of a supergame strategy. The model incorporates choices over strategies as well as beliefs about the strategies chosen by others. The strategically unsophisticated level-0 type randomizes uniformly over the set of strategies. The level-1 type holds beliefs within distance η of uniform and chooses a strategy within distance η of a best response to her beliefs, where $\eta \geq 0$ indexes the degree of error tolerance in terms of distance. The level-2 type holds beliefs within distance η of a distribution consistent with level-1 behavior, and again chooses a strategy within distance η of a best response to her beliefs. We measure distance in payoff space, defining distance for both choices and beliefs as proportional payoff loss, with belief payoffs given by the Quadratic Scoring Rule used in our experiment. We do not include higher level types in our estimation, but the model can be extended to include analogous higher level types.

Our second contribution is to show that our model of level- k reasoning can help to understand observed strategy choices and beliefs in repeated games. Specifically, using the IRPD dataset from Gill and Rosokha (2024), we classify each subject according to whether her strategy choices and her elicited beliefs are jointly consistent with the level-1 type, the level-2 type, both types, or neither type, at a given level of error tolerance. Classification therefore requires choosing the level of error tolerance, and this choice involves a trade-off. Tolerating too little error means that few subjects conform to the level- k model, while tolerating too much error means that the behavior of the level-1 and level-2 types overlaps excessively. Limiting error tolerance also guards against classifying as level-1 or level-2 those subjects whose behavior in fact arises from some other process. We resolve this trade-off by introducing the new notion of *balance-optimal error tolerance*, defined as the tolerance that maximizes the share of subjects classified uniquely as level-1 or level-2. At

⁵Formally, strategy elicitation induces a one-shot normal-form game whose set of Nash equilibria does not coincide exactly with the set of SPE of the direct-response IRPD (differences are due to the number of strategies and sequential rationality).

the balance-optimal error tolerance, we find that a substantial fraction of subjects are classified as level-1 or level-2 types. Based on the type classification, we also investigate how level- k types map empirically to beliefs and strategy choices.

We are not aware of any prior paper that measures level- k reasoning that operates at the level of a supergame strategy in a repeated game. Wunder et al. (2009)’s purely theoretical exercise calculates higher-order best responses for a cognitive hierarchy model defined over strategy choice in the finitely repeated Prisoner’s Dilemma. Castagnetti and Proto (2020, Section 5.1) note that the positive relationship between anger and AD documented by Castagnetti et al. (2018) is consistent with anger increasing level-0 behavior, because a level- k model where level-0 uniformly randomizes over $\{\text{AD}, \text{G}, \text{TFT}\}$ predicts that only the level-0 type chooses AD in their IRPD. Finally, Aoyagi et al. (2024, Section V) use numerical computation based on exogenously given parameters to show that a cognitive hierarchy model defined over strategy choice can generate patterns consistent with their repeated Prisoner’s Dilemma data; however, they do not estimate proportions of level- k types (nor do they use data on beliefs in this part of the paper).^{6,7}

Furthermore, we are not aware of any prior level- k paper that elicits beliefs over strategies (in repeated or sequential games); nor are we aware of any paper that jointly uses elicited beliefs and choices to estimate the distribution of level- k types. In the context of repeated games with level- k estimated round by round, Danz et al. (2012) estimate the distribution of level- k using only elicited beliefs, and separately using only choices, finding a high degree of within-individual consistency in estimated levels. In one-shot games, Hyndman et al. (2022) calculate the frequencies of different levels, averaged across games, again using elicited beliefs and choices separately.⁸

⁶A small literature measures level- k reasoning that operates at the level of a supergame strategy in the sequential Centipede game. García-Pola et al. (2020) use elicited strategies but not beliefs, while Kawagoe and Takizawa (2012) and Ho and Su (2013) elicit neither strategies nor beliefs. See also Schipper and Zhou (2024), who develop their model of strong level- k thinking for extensive-form games, and Lin and Palfrey (2024), who develop a dynamic cognitive hierarchy model for extensive-form games.

⁷Han et al. (2026) find that high-cognitive-skill gas station managers set higher prices. Using level- k estimated from a one-shot game (a variant of the 11-20 game), and using responses to unincentivized survey questions, they interpret gas station manager behavior through the lens of a level- k model of repeated pricing in which level-2 managers anticipate responses of rivals to price cuts, while level-1 managers do not.

⁸Other papers that use elicited beliefs in one-shot games in the context of level- k include Haruvy (2002), Costa-Gomes and Weizsäcker (2008), Rey-Biel (2009), Hyndman et al. (2012), Burchardi and Penczynski (2014), Lahav (2015), Goeree et al. (2018), Polonio and Coricelli (2019), Fragiadakis et al. (2020), Alempaki et al. (2022), Goeree and Garcia-Pola (2023), and Friedman and Ward (2025).

A growing literature studies the role of cognitive ability in strategic behavior (see the review by Sofianos, 2025). Cognitive ability supports logical reasoning and the organization of new information. In strategic settings, higher cognitive ability enhances strategic skill by improving individuals’ understanding of the strategic environment, which in turn facilitates the formation of beliefs about others and supports better choices given those beliefs (Fe et al., 2022). We measured cognitive ability using matrix reasoning, which captures fluid intelligence (i.e., logical or analytical intelligence) and is a leading measure of cognitive ability that correlates strongly with general intelligence; our test is similar to Raven’s Progressive Matrices (see Section II for details).

Our third contribution is to provide evidence that cognitive ability predicts higher level- k reasoning when level- k reasoning operates at the level of a supergame strategy. Specifically, using our IRPD dataset, we show that subjects with higher cognitive ability are more likely to be classified as level-2 rather than level-1. Our findings complement Proto et al. (2019, 2022), who find that cognitive ability influences cooperation in the IRPD because more cognitively able subjects make fewer errors when implementing their repeated-game strategies. Our design with strategy elicitation shuts down their primary mechanism.⁹ Instead, we find a novel channel through which cognitive ability influences behavior in the IRPD, namely through higher level reasoning in belief formation and strategy selection. The IRPD captures well complex strategy choice in repeated games that allow cooperation, punishment and forgiveness. To the extent that our findings extend to more general repeated games, accounting for higher level- k reasoning will sharpen predictions about how cognitive ability shapes behavior in those environments. We also complement the literature showing that cognitive ability predicts higher level- k reasoning that operates at the level of a single round (see Brañas-Garza et al., 2012, Gill and Prowse, 2016, Alós-Ferrer and Buckenmaier, 2021, Jin, 2021, Fe et al., 2022, and Ballester et al., 2024; although Georganas et al., 2015, find mostly null results).

Finally, we provide evidence that, when level- k reasoning operates at the level of a supergame strategy, subjects shift toward higher level- k reasoning as they gain experience. Specifically, subjects are more likely to be classified as level-2, and less likely as level-1, in the final supergame compared to the first. This finding complements the literature that documents shifts toward higher level reasoning with experience when level- k reasoning

⁹Consistent with this, Gill and Rosokha (2024, Appendix VIII) report that cognitive ability does not predict cooperation in our data. Gill and Rosokha (2024, Appendix VIII) also report that cognitive ability modestly predicts belief accuracy in the final supergame.

operates at the level of a single round (e.g., Stahl, 1996, Duffy and Nagel, 1997, Danz et al., 2012, Gill and Prowse, 2016, Ferraz et al., 2025).

The paper proceeds as follows. Section II describes the experimental design. Section III motivates the analysis by using a clustering algorithm to group subjects according to the similarity of their beliefs. Section IV develops our model of level- k reasoning in repeated games. Section V classifies subjects into levels of reasoning using the new notion of balance-optimal error tolerance. Section VI uses regression analysis to study the effects of cognitive ability, personality, and experience on level- k reasoning. Section VII widens the type classification and investigates how level- k types map empirically to beliefs and strategy choices. Section VIII shows that our results are robust to the choice of error tolerance. Section IX concludes, while the Appendix provides further details and experimental screenshots.

II Experimental design

The data come from our earlier study, Gill and Rosokha (2024). Here, we outline the experimental design (for further details and screenshots, see Section 2 and Appendices I and III of Gill and Rosokha, 2024).

We ran 27 sessions with 394 student subjects at Purdue University’s Vernon Smith Experimental Economics Laboratory. Subjects played 25 indefinitely repeated Prisoner’s Dilemma supergames with random rematching. The stage-game payoff matrix (Figure 1) follows Dal Bó and Fréchette (2011), with the return to joint cooperation $R \in \{32, 40, 48\}$, varied between subjects, and continuation probability $\delta = 0.75$.

	Cooperate	Defect
Cooperate	R, R	12, 50
Defect	50, 12	25, 25

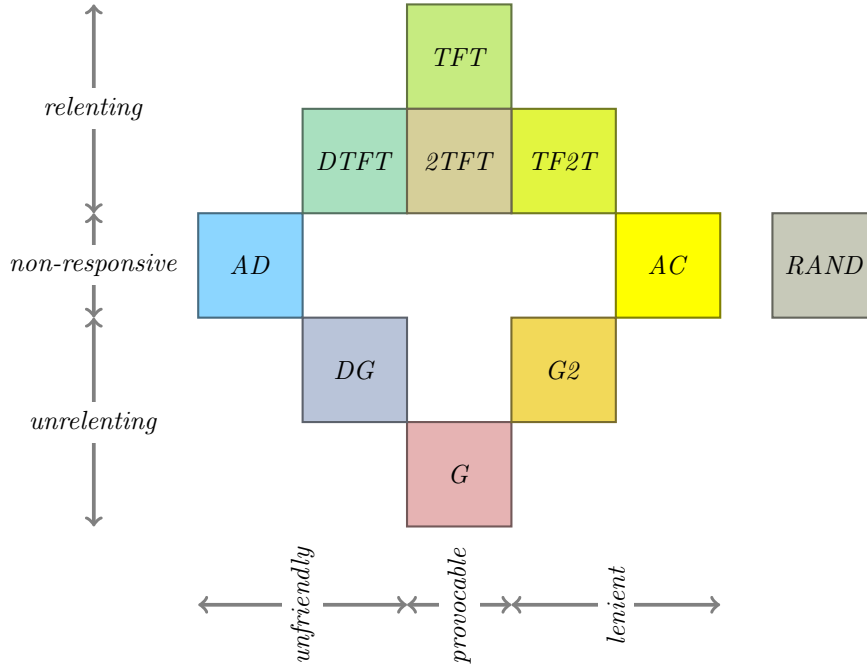
Notes: The experiment used labels ‘A’ and ‘B’ to represent Cooperate and Defect.

Figure 1: Stage-game payoff matrix

We elicited supergame strategies. At the start of each supergame, each subject chose one of 10 strategies to play the supergame, as described in Figure 2. These strategies include the 8 most common identified via structural estimation using round-by-round choices in the IRPD (Gill and Rosokha, 2024, Section I.D), together with DG (the unrelenting analog of DTFT) and RAND (which randomizes each round). The subject’s chosen strategy and that of her opponent were then played out round-by-round on the subject’s screen. The supergame history (choices and payoffs) was displayed during the supergame and after the supergame ended (until the subject chose to continue to the next supergame). To preserve external validity, subjects were not told the strategy chosen by their opponent. Appendix B, page 1, provides an illustrative screenshot.

In Figure 2, the horizontal axis categorizes the strategies according to when they first defect: ‘unfriendly’ strategies defect in the 1st round, while ‘provocable’ (‘lenient’) strategies defect (do not defect) immediately in response to the opponent’s first defection. The vertical axis categorizes the strategies according to whether, after punishing a rival’s

defection, they eventually relent and cooperate if the opponent cooperates.



Notes: AD: Always Defect; AC: Always Cooperate; G: Grim; G2: Lenient Grim 2; DG: Defect 1st round, then play G (ignoring own 1st-round defection); TFT: Tit-for-Tat; 2TFT: 2-Tits-for-1-Tat; TF2T: Tit-for-2-Tats; DTFT: Defect 1st round, then play TFT; RAND: Cooperate / Defect with probability $\frac{1}{2}$ in every round. Appendix B, page 1, provides the full wording of each strategy.

Figure 2: Description of the 10 strategies

To help subjects understand the 10 repeated-game strategies available in each of the 25 indefinitely repeated Prisoner’s Dilemma supergames, subjects completed two forms of training (without incentives) before we elicited strategies in the 1st supergame. First, each subject played 10 practice supergames against herself using the direct-response method (choosing an action for herself and for the ‘other’ in each round of each practice supergame). Second, subjects spent 3 minutes reading descriptions of the 10 available strategies, and then 5 minutes testing pairs of strategies against each other in training supergames played out round-by-round on the subject’s screen (on average, testing 16.2 pairs, including repeats). Figures A.3 to A.6 in Gill and Rosokha (2024) provide illustrative screenshots.

We also elicited beliefs about supergame strategies. We elicited beliefs twice, in the 1st supergame and the 25th (final) supergame. Specifically, we elicited beliefs about the

distribution of the 10 strategies, incentivized using the Quadratic Scoring Rule (QSR).¹⁰ To prevent contamination of initial strategy choice, in the 1st supergame we elicited beliefs unexpectedly after subjects had chosen their strategy (but before the supergame was played out). Appendix B, pages 2-3, provides illustrative screenshots.

We measured cognitive ability using matrix reasoning.¹¹ At the beginning of the experiment, subjects completed the 11-item test of matrix reasoning from the International Cognitive Ability Resource (Dworak et al., 2021; <https://icar-project.com>), which is similar to Raven’s Progressive Matrices (Raven et al., 2000), and which has been used in economics by, e.g., Gill et al. (2025). The test requires subjects to identify missing elements that complete visual patterns: Figure A.1 in Appendix A.1 provides a sample item. Matrix reasoning captures fluid intelligence (i.e., logical or analytical intelligence), and is a leading measure of cognitive ability that correlates strongly with general intelligence g (Gray and Thompson, 2004, Box 1) and with performance on other complex cognitive tasks (Carpenter et al., 1990). We standardize the matrix reasoning test scores, giving a cognitive ability measure with zero mean and unit standard deviation.

We measured personality using questions from the psychometric literature. At the beginning of the experiment, subjects completed a 52-item personality questionnaire. Based on the answers, and using principal factor analysis, we constructed five personality factors that measure: anxiety; cautiousness; kindness; manipulateness; and trust (for details see Gill and Rosokha, 2024, Appendix III.9). By construction, the personality factors are standardized (zero mean and unit standard deviation) and uncorrelated with each other. Furthermore, the personality factors all have correlations with cognitive ability that are smaller than 0.1.

Finally, we measured demographics at the end of the experiment. Specifically, we asked about age, gender, major, and whether the subject went to high school in the United States.

¹⁰We gave subjects no feedback about the accuracy of their guesses. Gill and Rosokha (2024, Appendix III.7) discuss properties of the QSR and evidence about effects of eliciting beliefs.

¹¹Matrix reasoning has been used to study how cognitive ability affects strategic behavior in a variety of games (e.g., Gill and Prowse, 2016, Proto et al., 2019, 2022, Fe et al., 2022, Hermes and Schunk, 2022).

III Belief clustering

Before developing our model of level- k reasoning in repeated games, we motivate our analysis by examining subjects' beliefs over strategies. To do so, we use a clustering algorithm to group subjects into clusters according to the similarity of their beliefs.

Specifically, we use affinity propagation clustering (Frey and Dueck, 2007), which is a popular unsupervised machine learning clustering algorithm with two key advantages. First, affinity propagation determines the optimal number of clusters endogenously. Second, within each cluster, affinity propagation selects one subject's beliefs to be the 'exemplar' beliefs that are representative of beliefs in that cluster. Affinity propagation has been used previously by, e.g., Romero and Rosokha (2018) to cluster repeated game strategies.

Panel (a) of Table 1 reports the output of the affinity propagation clustering algorithm for the beliefs over strategies elicited in the 1st supergame. Cluster 1, the largest cluster with 56 subjects, has exemplar beliefs that are exactly uniform. Clusters 2, 3, 5, 6 and 7 also have close-to-uniform beliefs (specifically, these clusters have exemplar beliefs whose mean-squared deviation from uniform does not exceed 25).

Together, the 187 subjects in these 6 clusters with close-to-uniform beliefs in the 1st supergame help to motivate the development of our model of level- k reasoning in repeated games, since uniform beliefs are consistent with a level-1 type who believes others are strategically unsophisticated (level-0) and therefore randomize uniformly over strategies. Our level- k analysis is further motivated by the finding that our subjects generally choose strategies that perform well given their beliefs: in the 1st supergame, around half of the subjects choose a strategy that achieves at least 95% of the expected payoff from the best response to their beliefs, while around a quarter of the subjects choose a strategy that is a perfect best response (panel (a) of Figure A.2 in Appendix A.1 provides the cumulative distribution function).

Cluster	N	<i>AD</i>	<i>DG</i>	<i>DTFT</i>	<i>RAND</i>	<i>G</i>	<i>2TFT</i>	<i>TFT</i>	<i>G2</i>	<i>TF2T</i>	<i>AC</i>
1	56	10	10	10	10	10	10	10	10	10	10
2	34	5	9	9	5	10	17	16	17	9	3
3	31	18	9	11	12	9	9	8	10	8	6
4	30	50	8	10	5	5	5	5	5	5	2
5	25	20	15	10	5	10	10	10	5	5	10
6	21	5	20	15	5	10	10	10	10	10	5
7	20	10	10	20	5	10	10	15	5	10	5
8	15	23	10	5	5	5	7	5	10	7	23
9	14	40	10	10	5	5	5	5	5	10	5
10	13	10	10	10	5	10	10	10	5	25	5
11	13	75	3	3	3	3	3	3	3	3	1
12	12	32	5	5	2	5	5	5	5	5	31
13	12	30	15	20	5	5	5	5	5	5	5
14	11	5	5	5	5	10	5	20	5	10	30
15	10	91	1	1	1	1	1	1	1	1	1
16	10	49	3	4	3	3	3	5	3	2	25

(a) Supergame 1

Cluster	N	<i>AD</i>	<i>DG</i>	<i>DTFT</i>	<i>RAND</i>	<i>G</i>	<i>2TFT</i>	<i>TFT</i>	<i>G2</i>	<i>TF2T</i>	<i>AC</i>
1	44	20	15	15	5	10	10	10	5	5	5
2	36	10	10	10	10	10	10	10	10	10	10
3	33	5	10	10	5	15	15	15	10	10	5
4	32	50	5	10	5	5	5	5	5	5	5
5	27	2	15	20	2	11	11	11	11	15	2
6	17	25	25	20	3	5	5	5	5	5	2
7	17	80	5	2	1	1	2	5	2	2	0
8	15	98	0	0	0	0	1	1	0	0	0
9	15	15	5	20	5	5	5	30	5	5	5
10	14	10	5	5	2	20	30	15	5	5	3
11	14	10	30	40	1	1	9	8	1	0	0
12	14	5	20	15	5	10	10	20	5	5	5
13	11	40	40	3	2	3	3	3	2	3	1
14	11	65	5	20	0	0	3	5	2	0	0
15	11	25	3	25	3	30	3	3	2	3	3
16	10	40	3	3	3	3	3	2	3	3	37

(b) Supergame 25

Notes: For each cluster, the table reports the exemplar beliefs. We report here clusters with $N \geq 10$; Tables A.1 and A.2 in Appendix A.1 report all the clusters (none of the smaller clusters has exemplar beliefs with mean-squared deviation from uniform ≤ 25). We implemented affinity propagation clustering in R using the APCluster package (Bodenhofer et al., 2011), and using default values for all arguments.

Table 1: Belief clusters

Panel (a) of Table 1 also shows other interesting patterns of beliefs. For example, a number of clusters have exemplar beliefs that place 50% or more weight on AD (e.g., clusters 4, 11 and 15), while others have exemplars that place 50% or more weight on AC (clusters 17, 22 and 23; see Table A.1 in Appendix A.1). An interesting belief pattern also emerges in which AD and AC both attract relatively high probability weights, with close-to-uniform beliefs among the remaining strategies (e.g., clusters 8, 12 and 16): these clusters suggest that some subjects' initial beliefs tend to gravitate toward simple memory-0 strategies.

Recall that we also elicited beliefs over strategies in the 25th (final) supergame: panel (b) of Table 1 reports the corresponding belief clusters. Clusters 1, 2 and 3 have close-to-uniform beliefs (again, these clusters have exemplar beliefs whose mean-squared deviation from uniform does not exceed 25), and we find 113 subjects in these clusters. The decline in the number of subjects in clusters with close-to-uniform beliefs, from 187 in the 1st supergame to 113 in the 25th supergame, suggests that some subjects moved from level-1 to higher level reasoning, which in turn motivates our later analysis of how experience affects level- k reasoning in Section VI.

IV Level- k model: Description

In this section, we develop a model of level- k reasoning in repeated games that incorporates choices over supergame strategies as well as beliefs about the supergame strategies chosen by others. We build on Nagel (1995)’s level- k model of boundedly rational reasoning, in which the level-0 type is strategically unsophisticated, the level-1 type best responds to the belief that all other players are level-0, the level-2 type best responds to the belief that all other players are level-1, and so on. As described by Crawford et al. (2013)’s review of the literature, level- k rules “anchor beliefs in an instinctive reaction to the game and then adjust them via a small number of iterated best responses.”¹²

The model describes behavior in one repeated game (i.e., in one supergame). We assume that players choose from a finite set of supergame strategies. We make this assumption for simplicity and to match our experimental data (described in Section II), noting also that evidence from strategy estimation based on round-by-round choices in the indefinitely repeated Prisoner’s Dilemma shows that most choices are consistent with a limited set of strategies (Dal Bó and Fréchette, 2018). We further assume that players form a belief about the supergame strategies chosen by others, in the form of a probability distribution over the finite set of strategies.

Assumption 0. *The level-0 type randomizes uniformly over the finite set of strategies.*

Uniform randomization by the strategically unsophisticated level-0 type is a common assumption in the level- k literature (e.g., Costa-Gomes et al., 2001, Crawford and Iriberri, 2007a, Fudenberg and Liang, 2019). Crawford et al. (2013)’s review notes that: “In most applications ... $L0$ is assumed to be uniform random over others’ feasible decisions...”

Assumption 1. *The level-1 type holds beliefs within distance η of uniform and chooses a strategy within distance η of a best response to her beliefs, where $\eta \in [0, 1]$ indexes the degree of error tolerance.*

As is standard in the level- k literature, we allow types to make errors (e.g., Ho et al., 1998, Gill and Prowse, 2016, and Ho et al., 2021, allow noise in choices, while Danz et al., 2012, allow noise in beliefs). Crawford et al. (2013)’s review notes that: “In empirical

¹²As further described by Nagel (1995), the level-1 type “forms first-order beliefs on the behavior of the other players. He thinks that others select a number at random, and he chooses his best response to this belief,” while the level-2 type “forms second-order beliefs on the first-order beliefs of the others...”

applications, it is assumed that $L1$ and higher types make errors...” In particular, $\eta \in [0, 1]$ indexes the degree of error tolerance in terms of distance, for both beliefs and choices. When $\eta = 0$, the error-free level-1 type believes that all other players are level-0, and so she holds uniform beliefs while choosing a best response to her uniform beliefs. When $\eta > 0$, the level-1 type’s beliefs can diverge from uniform by at most a distance of η , and the level-1 type’s strategy can diverge from a best response to her beliefs by, again, at most a distance of η .¹³

We use payoff space to measure distance. This choice aligns naturally with our experimental data, where both choices and beliefs were incentivized, but in other settings alternative distance metrics could be more appropriate. Specifically, we say that a player’s strategy s is distance $d_s \in [0, 1]$ from the best response to her beliefs if s achieves expected payoffs, given her beliefs, that are a proportion $1 - d_s$ of the expected payoffs from the best response to her beliefs.¹⁴ Furthermore, we say that a player’s beliefs b are distance $d_b \in [0, 1]$ from distribution x if reporting b when the true distribution is x achieves expected payoffs that are a proportion $1 - d_b$ of the expected payoffs from reporting x (as noted in Section II, beliefs were incentivized using the QSR). For comparability across the domains of choices and beliefs when using payoff space to measure distance, we normalize payoffs to range from 0 to 1 within each domain.¹⁵

Assumption 2. *The level-2 type holds beliefs that are within distance η of a distribution that places probability weight only on strategies that are consistent with level-1 behavior, and chooses a strategy within distance η of a best response to her beliefs.*

Just like for the level-1 type, $\eta \in [0, 1]$ indexes the degree of error tolerance in terms of distance, for the level-2 type’s beliefs and choices. We define “strategies that are consistent with level-1 behavior” as strategies that are within distance η of a best response to uniform beliefs. This definition implies that the level-2 type allows for potential error in the level-1 type’s choices. Previous models in the level- k literature that include higher level types who take into account noisy behavior by lower level types include Stahl and Wilson (1994),

¹³Noise in beliefs could stem from subjects whose beliefs deviate in a systematic way from the exact level- k belief prediction or whose beliefs deviate randomly every time they play the game. Either way, in our level- k model, players best respond (with noise) to their noisy belief. This is different from Costa-Gomes and Weizsäcker (2008), who assume that a latent belief drives behavior and is reported with noise. In our setting with learning and only two observations per subject, we cannot recover latent beliefs.

¹⁴With multiple best responses, this payoff-based distance is the same across all best responses.

¹⁵Alternatively, the model could be generalized to allow separate error tolerance parameters, one for choices and another for beliefs.

Ho et al. (1998), and Goeree et al. (2018).¹⁶ The definition further implies that the level-2 type assumes that the level-1 type holds uniform beliefs: thus, the level-2 type does not allow for error in the level-1 type’s beliefs. We make this assumption for conceptual and computational simplicity, noting that Danz et al. (2012) also assume a level- k type with noisy beliefs who does not allow for noise in the beliefs of the level- $(k - 1)$ type.

Although we do not include types higher than level-2 in our estimation (see Section V for discussion), the model could be extended to include higher types, with a level- k type (for $k > 2$) holding beliefs that are within distance η of a distribution that places probability weight only on strategies that are compatible with level- $(k - 1)$ behavior, while also choosing a strategy within distance η of a best response to her beliefs.

¹⁶As described by Crawford et al. (2013)’s review, the level-2 type in Stahl and Wilson (1994)’s model is “best responding to a noisy $L1$.” In Ho et al. (1998)’s model: “Level- L players are assumed to believe that all other players (besides themselves) choose from the Level $L - 1$ density,” where the lower level density incorporates noisy choice by the lower level agents. Goeree et al. (2018) use a level- k model with noisy introspection where “level-1 makes a noisy best response to level-0, level-2 makes a noisy best response to the noisy play of level-1, etc.”

V Level- k model: Estimation

V.A Estimation methodology

We follow, for example, Nyarko and Schotter (2002), Bellemare et al. (2008) and Spiliopoulos (2012) by using elicited beliefs, together with choices, to estimate a structural model of strategic behavior. We estimate our level- k model using the data from Gill and Rosokha (2024) described in Section II. Recall that we elicited beliefs about supergame strategies only in the 1st supergame and 25th (final) supergame, and so we estimate the model twice, first using the data from the 1st supergame, and second using the data from the 25th supergame.

We include only the level-1 and level-2 types in our estimation. We exclude higher types because in our setting with choices and beliefs over 10 strategies, and where subjects are given only the stage-game payoff matrix, the computational burden on higher types makes higher level behavior somewhat implausible. Furthermore, in mostly simpler settings, the level- k literature finds most weight on level-1 and level-2 behavior.¹⁷ We exclude the level-0 type because we cannot use beliefs to identify level-0 behavior. Furthermore, much of the level- k literature assumes that the level-0 type has zero frequency in the population, serving only to anchor the thinking of higher level types (e.g., Costa-Gomes and Crawford, 2006, Crawford and Iriberri, 2007b, Ho and Su, 2013). Crawford et al. (2013)’s review notes that “there is no presumption that $L0$ players exist: $L0$ is simply $L1$ ’s model of others...” and that “the estimated frequency of $L0$ is usually zero or small.”

We follow a common approach in the level- k literature by estimating the model subject by subject (e.g., Nagel, 1995, Costa-Gomes and Crawford, 2006, Alós-Ferrer and Buckenmaier, 2021). Recall that $\eta \in [0, 1]$ indexes error tolerance, for both choices and beliefs, with error tolerance increasing in η . We begin by estimating our level- k model for every

¹⁷Crawford and Iriberri (2007a) summarize the level- k literature as follows: “The estimated distribution tends to be stable across games, with most of the weight on $L1$ and $L2$.” For example, using first-round choices in the Beauty Contest game with $p < 1$, Duffy and Nagel (1997) find that level-1 and level-2 types make up 73% of the population, with only a very small proportion ($< 5\%$) of level-3 types.

value of η on a fine discrete grid.¹⁸ Specifically, for either Supergame 1 or Supergame 25, and for every value of η , we classify each subject into one of four mutually exclusive categories, as follows (Assumptions 1 and 2 in Section IV describe the level-1 and level-2 types).

Level- k classification:

- **L1:** Choices and beliefs jointly consistent only with the level-1 type.
- **L2:** Choices and beliefs jointly consistent only with the level-2 type.
- **Both:** Choices and beliefs jointly consistent with both the level-1 and level-2 types.
- **Neither:** Choices and beliefs jointly consistent with neither the level-1 nor level-2 type.

Note that we use the term ‘estimation’ loosely: for each value of error tolerance η , we simply classify each subject into a category. This approach allows us to classify subjects for any level of error tolerance, which in turn motivates our notion of balance-optimal error tolerance, as described below. We follow, e.g., Nagel (1995) in classifying subjects based on distance, while our approach contrasts with, e.g., Gill and Prowse (2016) who estimate a mixture-of-types level- k model of learning using maximum likelihood (Section V.D compares our classification methodology to maximum likelihood approaches).

V.B Results

The stacked graphs in Figure 3 present the classification results for the 1st supergame (panel (a)) and the 25th supergame (panel (b)), for values of error tolerance $\eta \leq 0.25$. When error tolerance is near zero, we classify few subjects as L1 or L2: unsurprisingly, few subjects conform to the level- k model when we allow near zero noise in choices or beliefs. As error tolerance increases from zero, at first we classify more and more subjects as L1 or L2. However, once error tolerance becomes large enough, the proportion of subjects classified as L1 or L2 begins to fall as η increases further. The reason is that when we

¹⁸We use the grid $\eta \in \{0, 0.00001, 0.00002, \dots, 1\}$. Payoff calculations are based on analytical calculations of expected payoffs for every pair of strategies, but are subject to computer precision. To establish whether a subject’s beliefs are consistent with the level-2 type, we use numerical optimization to minimize distance between the subject’s beliefs and distribution z , subject to the constraint that z places probability weight only on strategies that are consistent with level-1 behavior (specifically, we use the SLSQP constrained optimization algorithm from the `scipy.optimize` Python library, with multiple random initializations).

tolerate too much error, the level- k model loses predictive power: a wide range of behavior becomes consistent with both the level-1 type and the level-2 type, pushing many subjects to be classified as Both. This discussion motivates the notion of balance-optimal error tolerance.

Definition 1. *The balance-optimal error tolerance η^* is the value of η that maximizes the combined share of subjects classified as L1 or L2.*

That is, the balance-optimal error tolerance maximizes the share of subjects classified uniquely as level-1 or level-2. The balance-optimal error tolerance strikes a balance: tolerating enough error for the level- k model to explain observed behavior, while not tolerating so much error that the behavior of level-1 and level-2 types overlaps excessively. Next, we report the level- k classification at the balance-optimal error tolerance.

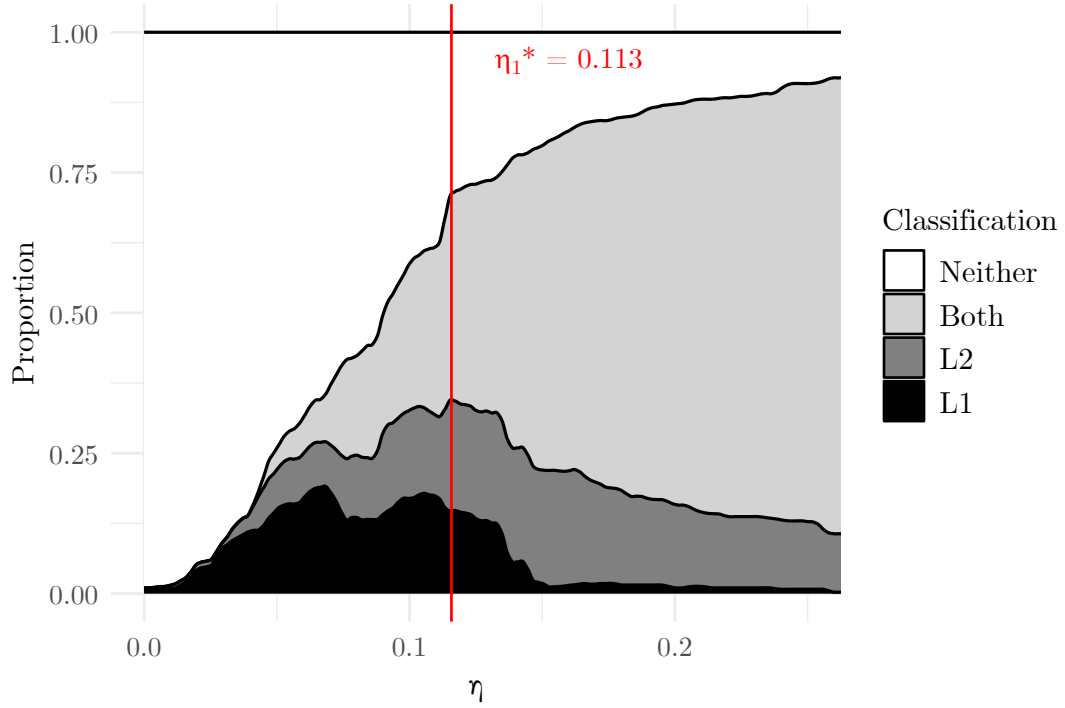
Result 1. *For the 1st supergame, the balance-optimal error tolerance $\eta_1^* = 0.113$, and at η_1^* we classify 14.7% of subjects as L1, 20.1% as L2, and 34.5% as Both.*

Result 2. *For the 25th supergame, the balance-optimal error tolerance $\eta_{25}^* = 0.118$, and at η_{25}^* we classify 7.6% of subjects as L1, 32.7% as L2, and 25.6% as Both.*

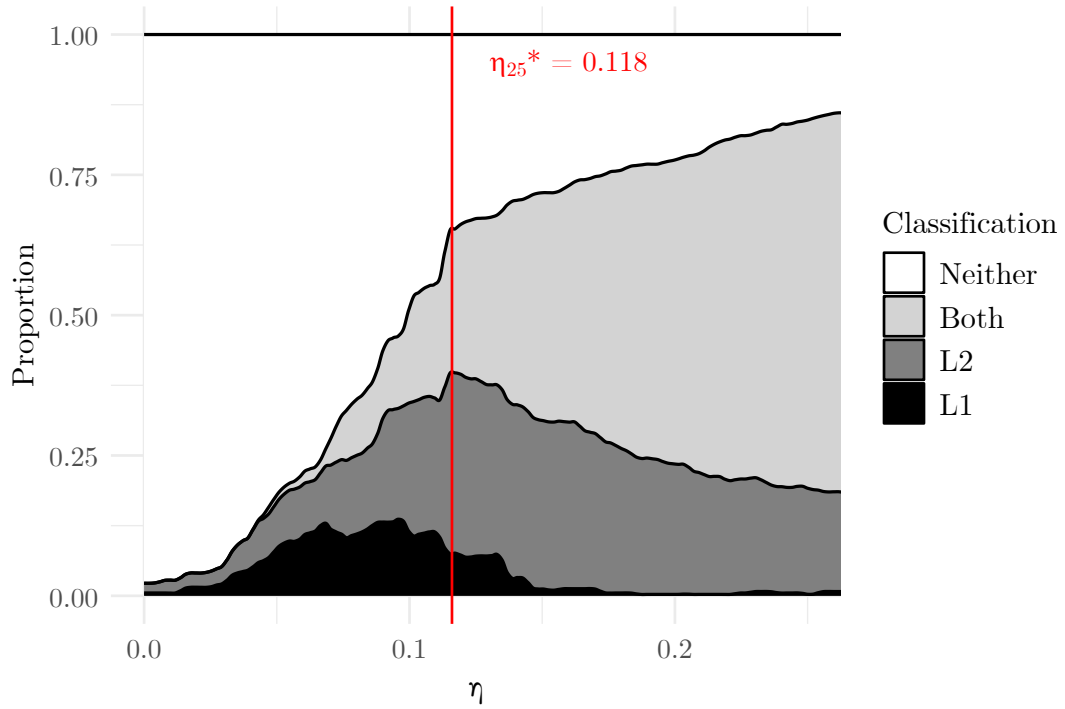
Comparing Result 1 with Result 2, and also comparing panels (a) and (b) of Figure 3 across the range of values of error tolerance, we can see that the level of reasoning increases with experience: as we move from the 1st supergame to the 25th (final) supergame, the proportion of subjects classified as L1 goes down, while the proportion classified as L2 goes up. In Section VI, we conduct econometric tests to confirm that this increase in level of reasoning with experience is statistically significant.

At the balance-optimal error tolerance we classify many subjects as Both. In Section VII, we introduce an alternative, wider level- k classification that reassigns these subjects to L1 or L2, using the fact that each of these subjects' beliefs are consistent with only the level-1 type or only the level-2 type when error tolerance η is lowered sufficiently. Section VII shows that our findings remain robust under this wider classification.

Finally, in Section VIII we show that our results are robust to the choice of error tolerance. First, to account for sampling variability, we bootstrap the balance-optimal error tolerance and re-run our analysis using error tolerances toward the ends of the bootstrapped range. Second, we classify all subjects as L1 or L2 without using the balance-optimal error tolerance at all.



(a) Supergame 1



(b) Supergame 25

Notes: The plots classify all 394 subjects and are smoothed using LOESS with a span of 0.01.

Figure 3: Level- k classification for range of values of error tolerance η

V.C Importance of beliefs

We emphasize that our level- k classification relies on using subjects' elicited beliefs, for two reasons. First, since beliefs consistent with different level- k types can rationalize the same strategy choice, we use beliefs to distinguish level- k types. For example, consider subject *q2hmdlzt* in the $R = 32$ treatment who chooses AD. We classify this subject as L2 because AD is a best response to the subject's level-2 beliefs that place 100% weight on AD. Without beliefs, we could not distinguish this subject from a level-1 type whose best response to uniform beliefs is also AD. We note that this difficulty in distinguishing level- k types without using elicited beliefs is particularly problematic in our complex setting. Given this complexity, we tolerate noise in beliefs (through the error tolerance parameter η), but this noise expands the set of strategies that can be rationalized by different level- k beliefs. Second, since the same strategy choice could arise from beliefs consistent with level- k reasoning or from some other process, we use beliefs to identify whether subjects are level- k types or not. For example, consider subject *yo91q5ok* in the $R = 32$ treatment who chooses AD. We classify this subject as Neither because her elicited beliefs are far from either level-1 or level-2 beliefs, placing 50% weight on AD and 50% weight on 2TFT. Without beliefs, and based only on choosing AD, we could not distinguish this subject from a level- k type.

V.D Further discussion of our methodology and connection to maximum likelihood

Our classification methodology has three advantages. First, it is direct and transparent. Fixing the error tolerance η , each subject's classification follows directly from her observed choices and beliefs, and we can observe clearly how the classification varies with the amount of error that we tolerate (Figures 3a and 3b). Second, error tolerance has a direct economic interpretation as the maximal proportional payoff loss allowed (for choices and for beliefs). Third, by separating level- k from non-level- k behavior in a structured data-driven way, our notion of balance-optimal error tolerance guards against classifying as level-1 or level-2 those subjects whose behavior in fact arises from some other process. Overall, our approach requires no parametric assumptions on the distribution of errors, no probabilistic allocation of subjects to types, and no arbitrary specification of non-level- k types.

By contrast, mixture-of-types maximum likelihood estimation is more complex, requiring the modeling of joint choice and belief distributions with assumptions on the error structure. Furthermore, allocating subjects to types requires using noisy posterior probabilities. Importantly, distinguishing level- k from non-level- k behavior within the maximum likelihood framework requires specifying particular non-level- k types. For example, Costa-Gomes et al. (2001) include dominance, equilibrium, and sophisticated types. In our approach, the Neither category instead accommodates non-level- k behavior without having to take a stand on what that behavior is.

We classify based on distance measured in payoff space. This contrasts with approaches that classify based on distance measured in the space of actions, such as the distance between chosen and level- k predicted guesses in the Beauty Contest (e.g., Nagel, 1995, Agranov et al., 2012, Alós-Ferrer and Buckenmaier, 2021). Since we measure distance in payoff space, our classification methodology has a close connection to maximum likelihood that directly classifies each individual (e.g., Costa-Gomes and Crawford, 2006), rather than estimating a mixture model. Noting that each subject’s choice is a fixed distance from the best response to her beliefs, and under the additional assumption that errors in choices and beliefs are independent, the maximum likelihood comparison of the level-1 and level-2 types reduces to comparing the distance in payoff space from the subject’s beliefs to each type’s level- k predicted beliefs (this holds true for any logit precision parameter). Thus, at any error tolerance η used to calculate the level- k predicted beliefs, and taking the subjects we classify as L1 or L2, maximum likelihood yields exactly the same classification.

VI Effects of cognitive ability and experience

In this section, we turn to regression analysis to study the effects of cognitive ability, personality, and experience on level- k reasoning. The regressions reported in Table 2 use the level- k classification from Section V reported in Results 1 and 2, and are based on data from the 1st and 25th supergames (the table notes provide full details).

In the 1st column, the dependent variable is an indicator for being classified as L2, and the sample includes only observations in which a subject is classified as either L1 or L2. Thus, the regression measures the effects of cognitive ability and personality on the probability of being classified as L2, conditional on being L1 or L2. The results show that a 1 standard deviation increase in cognitive ability is associated with a 9 percentage point increase in the probability of being classified as L2 ($p = 0.002$).

In the 2nd column, we further measure the effect of experience. The regression is the same as in the 1st column, except that we add an indicator for the 25th (final) supergame. The coefficient on this indicator shows that the probability of being classified as L2 in the 25th supergame is 22 percentage points higher than in the 1st supergame ($p < 0.001$).

In the columns headed “All Data”, we expand the sample to include all observations (that is, the sample now further includes observations in which a subject is classified as Both or Neither). The 5th and 6th columns (All Data, L2) show that cognitive ability continues to predict the probability of being classified as L2 ($p = 0.007$). In the 6th column, the coefficient on the 25th supergame shows that the probability of being classified as L2 continues to be higher with experience ($p = 0.001$). Finally, the 3rd and 4th columns (All Data, L1) show weaker, negative effects of cognitive ability and experience on the probability of being classified as L1.

We summarize the main findings from Table 2 as follows:

Result 3. *Cognitive ability predicts higher level- k reasoning.*

Result 4. *With experience, subjects shift toward higher level- k reasoning.*

Dep. Var.	L1 or L2 Subset		All Data		All Data	
	L2	L2	L1	L1	L2	L2
Cognitive Ability	0.090*** (0.026)	0.093*** (0.025)	−0.016* (0.009)	−0.016* (0.009)	0.046*** (0.015)	0.046*** (0.015)
Trust	−0.014 (0.034)	−0.015 (0.032)	−0.009 (0.015)	−0.009 (0.015)	−0.035** (0.017)	−0.035** (0.017)
Cautiousness	0.015 (0.033)	0.016 (0.032)	−0.007 (0.012)	−0.007 (0.012)	0.021 (0.018)	0.021 (0.018)
Anxiety	0.030 (0.027)	0.034 (0.027)	0.000 (0.011)	0.000 (0.011)	0.041** (0.020)	0.041** (0.020)
Kindness	0.024 (0.029)	0.026 (0.029)	0.003 (0.012)	0.003 (0.012)	0.021 (0.017)	0.021 (0.017)
Manipulativeness	0.014 (0.027)	0.017 (0.027)	−0.006 (0.015)	−0.006 (0.015)	0.006 (0.019)	0.006 (0.019)
Supergame 25		0.223*** (0.051)		−0.072** (0.028)		0.128*** (0.034)
Mean Dep. Var.	0.699	0.699	0.113	0.113	0.262	0.262
Num. obs.	292	292	780	780	780	780
N Clusters	27	27	27	27	27	27

Notes: Col. 1 reports an OLS regression where the dependent variable is an indicator for being classified as L2. The regression uses data from Supergame 1 and Supergame 25, and the sample includes only observations in which a subject is classified as either L1 or L2 at the balance-optimal error tolerance for the given supergame (subjects may therefore appear zero, one, or two times). Col. 2 adds an indicator for Supergame 25. Cols. 5 and 6 are the same as Cols. 1 and 2, except that the sample includes all observations. Cols. 3 and 4 are the same as Cols. 5 and 6, except that the dependent variable is an indicator for being classified as L1. Cognitive ability and personality measures are standardized. All regressions include controls for treatment and demographic characteristics (see Section II), specifically binary indicators for $R = 40$, $R = 48$, and for each characteristic. As in Gill and Rosokha (2024): (i) we binarize major into STEM / not STEM; (ii) we exclude from the regressions 4 subjects who answered “prefer not to say” to one or more demographic questions. Heteroskedasticity-robust standard errors clustered at the session level are in parentheses. ***, ** and * denote significance at the 1%, 5% and 10% levels (two-sided tests).

Table 2: Effects of cognitive ability, personality, and experience on level- k

We find no evidence of an interaction between the effects of cognitive ability and experience. When we add an interaction between cognitive ability and the indicator for the 25th supergame to the regressions reported in the 2nd, 4th and 6th columns of Table 2, in every case the coefficient on the interaction is far from statistical significance ($p > 0.7$).

Finally, we find little systematic evidence that personality traits predict level- k reasoning. We find some evidence that subjects higher in anxiety are more likely to be classified as L2: the effect of anxiety is significant at the 5% level in the 5th and 6th columns, and we find stronger effects of anxiety when we widen the classification of L1 and L2 (see Table 3 in Section VII).¹⁹ Although trust is statistically significant at the 5% level in the 5th and 6th columns, it is not significant in the other columns, and the effects of trust are no longer significant at the 5% level once we widen the classification of L1 and L2.²⁰

¹⁹Anxiety is one facet of neuroticism. Gill and Prowse (2016) find a positive association between a factor that loads on both agreeableness and emotional stability (the reverse of neuroticism) and higher level- k in a repeated Beauty Contest game.

²⁰Although not our focus of interest, our dataset includes gender as a demographic control: perhaps intriguingly, the coefficient on the gender control shows that male subjects are more likely to be classified as L2 in our data ($p < 0.01$ in the regressions reported in the 5th and 6th columns of Table 2). In one-shot games, Hyndman et al. (2022) categorize subjects into ‘stayers’ whose level- k beliefs are stable across games and ‘movers’ whose level- k beliefs change across games, finding that stayers have level-1 beliefs and are more likely to be female, while movers tend to have higher level beliefs and are more likely to be male (the results on gender are in Appendix C.3). Again in one-shot games, and using the cognitive hierarchy model, Camerer et al. (2004) also find higher level thinking among males. Brañas-Garza et al. (2012), however, find no gender difference in level- k .

VII Wider classification and behavior of L1 and L2

In this section, we widen the classification of L1 and L2 to include subjects that we previously classified as Both, and we show that our earlier results are robust to this wider classification. We also use the wider classification to investigate how L1 and L2 map empirically to beliefs and choices.

Section V classified each subject into one of four categories: L1, L2, Both, or Neither. Results 1 and 2 report the level- k classification: at the balance-optimal error tolerance for the given supergame, we classified many subjects as Both. Consider a subject classified as Both: this subject has choices and beliefs that are jointly consistent with both the level-1 and level-2 types. As we lower error tolerance η below the balance-optimal level η^* , eventually we reach a value at which the subject's beliefs are consistent with one and only one of the level-1 and level-2 types (Assumptions 1 and 2 define belief consistency). Here, we reassign the subject from Both to L1 (L2) if her beliefs are consistent only with the level-1 (level-2) type at this lower η . Using this methodology, we are able to reassign all subjects previously in the Both category as L1 or L2.

Based on this wider classification, for the 1st supergame we now classify 41.6% of subjects as L1 and 27.7% as L2, and for the 25th supergame we now classify 28.9% of subjects as L1 and 37.1% as L2. By construction, these percentages sum to the total of L1, L2 and Both from Results 1 and 2. Using the wider classification, we continue to find that the level of reasoning increases with experience: as we move from the 1st supergame to the 25th (final) supergame, the proportion of subjects classified as L1 continues to go down, while the proportion classified as L2 continues to go up.

With the wider classification of L1 and L2, we continue to find statistical support for Results 3 and 4 from Section VI. Table 3 reports the same regressions as in Table 2, but now using the wider classification. In particular, the positive effects of cognitive ability and experience on the probability of being classified as L2 continue to be statistically significant, with $p < 0.05$ in all cases. As previously noted in Section VI, we also find that subjects higher in anxiety are more likely to be classified as L2.

Dep. Var.	L1 or L2 Subset		All Data		All Data	
	L2	L2	L1	L1	L2	L2
Cognitive Ability	0.076*** (0.027)	0.073*** (0.026)	-0.038* (0.020)	-0.038* (0.020)	0.052*** (0.018)	0.052*** (0.018)
Trust	-0.027 (0.023)	-0.028 (0.023)	0.002 (0.019)	0.002 (0.019)	-0.034* (0.017)	-0.034* (0.017)
Cautiousness	0.018 (0.026)	0.019 (0.026)	-0.009 (0.019)	-0.009 (0.019)	0.024 (0.021)	0.024 (0.021)
Anxiety	0.056*** (0.020)	0.057*** (0.020)	-0.034** (0.016)	-0.034** (0.016)	0.044** (0.017)	0.044** (0.017)
Kindness	0.020 (0.024)	0.019 (0.024)	-0.012 (0.020)	-0.012 (0.020)	0.014 (0.019)	0.014 (0.019)
Manipulativeness	0.016 (0.021)	0.016 (0.021)	-0.005 (0.017)	-0.005 (0.017)	0.011 (0.018)	0.011 (0.018)
Supergame 25		0.160*** (0.032)		-0.131*** (0.028)		0.095** (0.036)
Mean Dep. Var.	0.475	0.475	0.355	0.355	0.322	0.322
Num. obs.	528	528	780	780	780	780
N Clusters	27	27	27	27	27	27

Notes: See the notes to Table 2. Here, we run the same regressions, except that we use the wider classification of L1 and L2 described in the main text (which reassigns subjects in Both to L1 or L2). Heteroskedasticity-robust standard errors clustered at the session level are in parentheses. ***, ** and * denote significance at the 1%, 5% and 10% levels (two-sided tests).

Table 3: Effects of cognitive ability, personality, and experience on level- k :
Wider classification of L1 and L2

Next, we use the wider classification to investigate how L1 and L2 map empirically to beliefs and choices, split by treatment. Specifically, Table 4 reports the empirical beliefs and strategy choices that we observe in the data for each treatment-category pair. We use the wider classification for this purpose because splitting by treatment with four categories (L1, L2, Both, and Neither) cuts the data too finely, with some treatment-category pairs containing few subjects. Table 4 reports data for the 1st supergame: Table A.3 in Appendix A.1 shows that the patterns described below persist in the 25th supergame.

Consistent with our level- k model, Table 4 shows that the average belief weights of L1 subjects are approximately uniform.²¹ When the return to joint cooperation is low ($R = 32$), L1 subjects favor unfriendly strategies that defect in the 1st round (AD, DTFT and DG, which are the three highest-expected-payoff strategies given uniform beliefs when $R = 32$). When the return to joint cooperation is high ($R = 48$), L1 subjects shift toward strategies that cooperate in the 1st round, with G2 being the most popular (G2 is the unique best response given uniform beliefs when $R = 48$).

When the return to joint cooperation is low ($R = 32$), Table 4 shows that L2 subjects place most belief weight on unfriendly strategies, broadly matching L1 behavior (L1 choose unfriendly strategies 76% of the time, while L2's beliefs place 72% weight on unfriendly strategies, skewed toward AD). Consistent with this pessimism, L2 choose AD 84% of the time. When the return to joint cooperation is high ($R = 48$), L2 beliefs become substantially more optimistic, in line with L1's more cooperative strategy choices. Consistent with this optimism, L2 subjects never choose unfriendly strategies.

Finally, Table 4 provides evidence that the Neither category includes some strategically unsophisticated subjects. First, the average beliefs of subjects classified as Neither place slightly *more* weight on unfriendly strategies when the return to joint cooperation is high ($R = 48$) than when it is low ($R = 32$), moving in the opposite direction to incentives. Second, only subjects classified as Neither choose AC when $R = 32$ or AD when $R = 48$. Third, they choose RAND more frequently than L1 or L2 subjects.

²¹Given our level- k model and classification methodology, this pattern is not surprising.

Classification	AD	DG	DTFT	RAND	G	2TFT	TFT	G2	TF2T	AC
R = 32										
L1	14	12	12	7	10	11	10	9	8	6
L2	56	10	6	4	4	3	4	4	3	5
Neither	19	10	9	7	8	10	11	7	8	12
R = 40										
L1	14	9	10	7	10	11	10	9	9	11
L2	34	6	6	3	7	7	11	6	8	12
Neither	14	13	12	5	12	9	9	8	7	11
R = 48										
L1	12	9	9	7	11	11	11	10	10	10
L2	7	4	3	3	14	11	12	9	9	29
Neither	26	11	12	5	9	7	8	7	6	8

(a) Beliefs in Supergame 1 (%)

Classification	AD	DG	DTFT	RAND	G	2TFT	TFT	G2	TF2T	AC
R = 32										
L1	27	29	20	0	5	8	8	0	2	0
L2	84	12	3	0	0	0	0	0	0	0
Neither	9	12	12	14	7	14	9	7	5	12
R = 40										
L1	26	2	4	0	9	25	9	9	8	8
L2	45	6	2	0	8	6	12	6	10	4
Neither	38	17	21	21	0	0	0	0	0	4
R = 48										
L1	0	0	2	0	17	15	13	27	10	15
L2	0	0	0	0	11	14	7	18	14	36
Neither	39	19	15	9	7	4	6	0	0	2

(b) Strategy choices in Supergame 1 (%)

Notes: R is the return to joint cooperation, varied between subjects (see Section II). We classify subjects using the wider classification described in the main text. For each treatment-category pair, we report the mean belief weight placed on each strategy (panel (a)) and the percentage of subjects choosing each strategy (panel (b)).

Table 4: Empirical beliefs and strategies by treatment-category pair in Supergame 1

VIII Robustness: Choice of error tolerance

Our notion of balance-optimal error tolerance maximizes the share of subjects classified uniquely as level-1 or level-2, and therefore strikes a balance by tolerating enough error for the level- k model to explain observed behavior, without tolerating so much error that the behavior of different types overlaps too much. A natural question remains: are our results sensitive to using the balance-optimal error tolerance?

We answer this question in two ways. First, we show that our results are robust when we allow for sampling variability by bootstrapping the balance-optimal error tolerance and then using error tolerances toward the ends of the range of bootstrapped values. In particular, Table 5 shows that the regression results in Table 2 are robust when we use error tolerances at the 5th percentile or 95th percentile of the distribution of bootstrapped balance-optimal error tolerances (the table notes provide details). Cognitive ability continues to predict higher level- k reasoning, and the level of reasoning continues to increase with experience. Specifically, the positive effects of cognitive ability and experience on the probability of being classified as L2 continue to be statistically significant, with $p < 0.01$ in all cases.²² Similarly to Table 2, we continue to find little systematic evidence of an effect of personality traits. Thus, our conclusions do not hinge on the exact value of the balance-optimal error tolerance.

Second, we show that our results are robust when we do not use the balance-optimal error tolerance at all. In particular, Table 6 shows that the regression results in Table 3 are robust when we use the wider classification methodology from Section VII, but with $\eta = 1$ instead of the balance-optimal error tolerance (the notes to Table 6 provide details). At $\eta = 1$, any amount of noise is tolerated, and thus we classify all subjects as L1 or L2. We find that the positive effects of cognitive ability and experience on the probability of being classified as L2 continue to be statistically significant, with $p < 0.05$ in all cases.²³

²²Compared to Table 2, the effect of cognitive ability on the probability of being classified as L2 in Cols. 1-2 of Table 5 is larger at the 5th percentile and smaller at the 95th percentile (although effects are similar in Cols. 5-6). We note that higher error tolerance admits more noise relative to the exact level- k predictions, potentially making cognitive ability less predictive for the L1 vs. L2 classification.

²³Compared to Table 3, the effect of cognitive ability on the probability of being classified as L2 is smaller in magnitude: again, as noted above, higher error tolerance admits more noise in level- k behavior, potentially making cognitive ability less predictive.

Dep. Var.	L1 or L2 Subset		All Data		All Data	
	L2	L2	L1	L1	L2	L2
Cognitive Ability	0.139*** (0.037)	0.131*** (0.035)	-0.032** (0.015)	-0.032** (0.015)	0.044*** (0.013)	0.044*** (0.013)
Trust	0.002 (0.036)	0.007 (0.036)	-0.007 (0.012)	-0.007 (0.012)	0.001 (0.015)	0.001 (0.015)
Cautiousness	0.008 (0.028)	0.006 (0.027)	-0.001 (0.007)	-0.001 (0.007)	0.014 (0.019)	0.014 (0.019)
Anxiety	0.044 (0.027)	0.040 (0.026)	-0.006 (0.012)	-0.006 (0.012)	0.026 (0.017)	0.026 (0.017)
Kindness	0.038 (0.031)	0.039 (0.030)	-0.008 (0.013)	-0.008 (0.013)	0.014 (0.016)	0.014 (0.016)
Manipulativeness	0.025 (0.032)	0.025 (0.032)	-0.005 (0.016)	-0.005 (0.016)	0.009 (0.016)	0.009 (0.016)
Supergame 25		0.235*** (0.053)		-0.067** (0.031)		0.090*** (0.031)
Mean Dep. Var.	0.582	0.582	0.141	0.141	0.196	0.196
Num. obs.	263	263	780	780	780	780
N Clusters	27	27	27	27	27	27

(a) At 5th percentile of distribution of bootstrapped balance-optimal error tolerances

Dep. Var.	L1 or L2 Subset		All Data		All Data	
	L2	L2	L1	L1	L2	L2
Cognitive Ability	0.071*** (0.024)	0.077*** (0.023)	-0.015* (0.009)	-0.015* (0.009)	0.043*** (0.015)	0.043*** (0.015)
Trust	-0.008 (0.036)	-0.015 (0.034)	-0.013 (0.013)	-0.013 (0.013)	-0.039** (0.016)	-0.039** (0.016)
Cautiousness	0.032 (0.033)	0.034 (0.031)	-0.011 (0.012)	-0.011 (0.012)	0.025 (0.018)	0.025 (0.018)
Anxiety	0.022 (0.029)	0.030 (0.029)	-0.002 (0.011)	-0.002 (0.011)	0.033 (0.020)	0.033 (0.020)
Kindness	0.037 (0.026)	0.041 (0.027)	-0.009 (0.010)	-0.009 (0.010)	0.019 (0.017)	0.019 (0.017)
Manipulativeness	-0.008 (0.025)	-0.007 (0.025)	0.004 (0.012)	0.004 (0.012)	0.002 (0.020)	0.002 (0.020)
Supergame 25		0.215*** (0.050)		-0.062** (0.023)		0.123*** (0.032)
Mean Dep. Var.	0.712	0.712	0.103	0.103	0.254	0.254
Num. obs.	278	278	780	780	780	780
N Clusters	26	26	27	27	27	27

(b) At 95th percentile of distribution of bootstrapped balance-optimal error tolerances

Notes: See the notes to Table 2. Here, we run the same regressions, except that we replace balance-optimal error tolerance values with percentiles from the distribution of bootstrapped balance-optimal error tolerances. In panel (a), we replace $\eta_1^* = 0.113$ and $\eta_{25}^* = 0.118$ with 0.098 and 0.100, the values at the 5th percentile for Supergame 1 and 25, respectively. In panel (b), we replace $\eta_1^* = 0.113$ and $\eta_{25}^* = 0.118$ with 0.128 and 0.124, the values at the 95th percentile for Supergame 1 and 25, respectively. We bootstrap at the session level, stratified by treatment (value of R), with 10,000 replications.

Table 5: Effects of cognitive ability, personality, and experience on level- k :
Robustness using bootstrapped balance-optimal error tolerance

Dep. Var.	L1 or L2 Subset		All Data		All Data	
	L2	L2	L1	L1	L2	L2
Cognitive Ability	0.045** (0.020)	0.045** (0.020)	-0.045** (0.020)	-0.045** (0.020)	0.045** (0.020)	0.045** (0.020)
Trust	-0.027 (0.021)	-0.027 (0.021)	0.027 (0.021)	0.027 (0.021)	-0.027 (0.021)	-0.027 (0.021)
Cautiousness	0.005 (0.024)	0.005 (0.024)	-0.005 (0.024)	-0.005 (0.024)	0.005 (0.024)	0.005 (0.024)
Anxiety	0.033* (0.018)	0.033* (0.018)	-0.033* (0.018)	-0.033* (0.018)	0.033* (0.018)	0.033* (0.018)
Kindness	0.022 (0.020)	0.022 (0.020)	-0.022 (0.020)	-0.022 (0.020)	0.022 (0.020)	0.022 (0.020)
Manipulativeness	0.015 (0.017)	0.015 (0.017)	-0.015 (0.017)	-0.015 (0.017)	0.015 (0.017)	0.015 (0.017)
Supergame 25		0.136*** (0.032)		-0.136*** (0.032)		0.136*** (0.032)
Mean Dep. Var.	0.422	0.422	0.578	0.578	0.422	0.422
Num. obs.	780	780	780	780	780	780
N Clusters	27	27	27	27	27	27

Notes: See the notes to Table 3. Here, we run the same regressions, using the same wider classification of L1 and L2 methodology, but setting error tolerance at its maximum value of $\eta = 1$ instead of at balance-optimal η_1^* and η_{25}^* . Specifically, at $\eta = 1$, all subjects are classified as Both. Consider a given subject in either Supergame 1 or 25: as we lower error tolerance below $\eta = 1$, eventually we reach a value at which the subject's beliefs are consistent with one and only one of the level-1 and level-2 types (Assumptions 1 and 2 define belief consistency). Here, we assign the subject to L1 (L2) if her beliefs are consistent only with the level-1 (level-2) type at this lower η . Using this methodology, we are able to classify all subjects as L1 or L2. Therefore: (i) the L1-or-L2 subset coincides with the full sample, so Cols. 1-2 replicate Cols. 5-6; and (ii) L1 effects in Cols. 3-4 are the negatives of L2 effects in Cols. 5-6. However, we keep all six columns for consistency and comparison with earlier tables.

Table 6: Effects of cognitive ability, personality, and experience on level- k :
Robustness without using balance-optimal error tolerance

IX Conclusion

This paper advances our understanding of boundedly rational behavior in repeated strategic interactions by introducing a novel model of level- k reasoning that operates at the level of supergame strategies. Unlike existing approaches that analyze reasoning round by round, our framework captures how individuals form and act on beliefs about entire strategic plans over the course of a repeated game. Using experimental data from the Indefinitely Repeated Prisoner’s Dilemma, where both strategies and beliefs were elicited, we demonstrate that many subjects exhibit behavior consistent with level-1 or level-2 reasoning. Importantly, because the level- k model is fundamentally belief-driven, the joint elicitation of strategies and beliefs facilitates identification of distinct reasoning processes. We find systematic variation in the depth of reasoning: individuals with greater cognitive ability and more experience are more likely to be classified as higher level types. These findings contribute to the behavioral foundations of repeated game theory and highlight the importance of belief formation and individual heterogeneity in strategic environments that unfold over time.

Our work provides a foundation for several promising directions for future research. First, our model of level- k reasoning could be applied to additional environments, including repeated coordination games, bargaining games, and market entry games. Second, the model could be extended to include more behavioral types, such as worldly players who best respond to the empirical distribution of behavior or types that incorporate social preferences. Third, the model of level- k reasoning could incorporate explicit belief updating, in order to investigate how players learn about the distribution of types over time and how this learning relates to individual characteristics. Finally, incorporating heterogeneous cognitive costs or bounded memory would enable richer predictions about when and why players adopt different levels of reasoning in repeated games.

References

- Agranov, M., Potamites, E., Schotter, A., and Tergiman, C. (2012). Beliefs and endogenous cognitive levels: An experimental study. *Games and Economic Behavior*, 75(2): 449–463
- Alaoui, L. and Penta, A. (2016). Endogenous depth of reasoning. *Review of Economic Studies*, 83(4): 1297–1333
- Alaoui, L. and Penta, A. (2022). Cost-benefit analysis in reasoning. *Journal of Political Economy*, 130(4): 881–925
- Alempaki, D., Colman, A.M., Kölle, F., Loomes, G., and Pulford, B.D. (2022). Investigating the failure to best respond in experimental games. *Experimental Economics*, 25(2): 656–679
- Alós-Ferrer, C. and Buckenmaier, J. (2021). Cognitive sophistication and deliberation times. *Experimental Economics*, 24(2): 558–592
- Aoyagi, M., Fréchette, G.R., and Yuksel, S. (2024). Beliefs in repeated games: An experiment. *American Economic Review*, 114(12): 3944–3975
- Arad, A. and Rubinstein, A. (2012). The 11–20 money request game: A level- k reasoning study. *American Economic Review*, 102(7): 3561–3573
- Ballester, C., Rodriguez-Moral, A., and Vorsatz, M. (2024). Cognitive reflection in experimental anchored guessing games. *Games and Economic Behavior*, 148: 179–195
- Bellemare, C., Kröger, S., and Van Soest, A. (2008). Measuring inequity aversion in a heterogeneous population using experimental decisions and subjective probabilities. *Econometrica*, 76(4): 815–839
- Benndorf, V., Kübler, D., and Normann, H.T. (2017). Depth of reasoning and information revelation: An experiment on the distribution of k -levels. *International Game Theory Review*, 19(04): 1750021
- Bodenhofer, U., Kothmeier, A., and Hochreiter, S. (2011). APCluster: an R package for affinity propagation clustering. *Bioinformatics*, 27(17): 2463–2464
- Brañas-Garza, P., García-Muñoz, T., and González, R.H. (2012). Cognitive effort in the beauty contest game. *Journal of Economic Behavior & Organization*, 83(2): 254–260
- Brandts, J. and Charness, G. (2011). The strategy versus the direct-response method: a first survey of experimental comparisons. *Experimental Economics*, 14(3): 375–398
- Burchardi, K.B. and Penczynski, S.P. (2014). Out of your mind: Eliciting individual reasoning in one shot games. *Games and Economic Behavior*, 84: 39–57
- Camerer, C.F., Ho, T.H., and Chong, J.K. (2004). A cognitive hierarchy model of games. *Quarterly Journal of Economics*, 119(3): 861–898

- Carpenter, P.A., Just, M.A., and Shell, P.** (1990). What one intelligence test measures: A theoretical account of the processing in the Raven Progressive Matrices test. *Psychological Review*, 97(3): 404–413
- Cason, T.N. and Mui, V.L.** (2019). Individual versus group choices of repeated game strategies: A strategy method approach. *Games and Economic Behavior*, 114: 128–145
- Castagnetti, A. and Proto, E.** (2020). Anger and strategic behavior: A level- k analysis. *IZA Discussion Paper 13661*
- Castagnetti, A., Proto, E., and Sofianos, A.** (2023). Anger impairs strategic behavior: A Beauty-Contest based analysis. *Journal of Economic Behavior & Organization*, 213: 128–141
- Castagnetti, S.A., Massaro, S., and Proto, E.** (2018). The influence of anger on strategic cooperative interactions. *Proceedings of the Academy of Management Annual Meeting*, 14162
- Costa-Gomes, M., Crawford, V.P., and Broseta, B.** (2001). Cognition and behavior in normal-form games: An experimental study. *Econometrica*, 69(5): 1193–1235
- Costa-Gomes, M.A. and Crawford, V.P.** (2006). Cognition and behavior in two-person guessing games: An experimental study. *American Economic Review*, 96(5): 1737–1768
- Costa-Gomes, M.A. and Weizsäcker, G.** (2008). Stated beliefs and play in normal-form games. *Review of Economic Studies*, 75(3): 729–762
- Crawford, V.P. and Iriberri, N.** (2007a). Level- k auctions: Can a nonequilibrium model of strategic thinking explain the winner’s curse and overbidding in private-value auctions? *Econometrica*, 75(6): 1721–1770
- Crawford, V.P., Costa-Gomes, M.A., and Iriberri, N.** (2013). Structural models of nonequilibrium strategic thinking: Theory, evidence, and applications. *Journal of Economic Literature*, 51(1): 5–62
- Crawford, V.P. and Iriberri, N.** (2007b). Fatal attraction: Salience, naïveté, and sophistication in experimental “Hide-and-Seek” games. *American Economic Review*, 97(5): 1731–1750
- Dal Bó, P. and Fréchette, G.R.** (2011). The evolution of cooperation in infinitely repeated games: Experimental evidence. *American Economic Review*, 101(1): 411–29
- Dal Bó, P. and Fréchette, G.R.** (2018). On the determinants of cooperation in infinitely repeated games: A survey. *Journal of Economic Literature*, 56(1): 60–114
- Dal Bó, P. and Fréchette, G.R.** (2019). Strategy choice in the infinitely repeated prisoner’s dilemma. *American Economic Review*, 109(11): 3929–52
- Danz, D.N., Fehr, D., and Kübler, D.** (2012). Information and beliefs in a repeated normal-form game. *Experimental Economics*, 15(4): 622–640

- Duffy, J.** and **Nagel, R.** (1997). On the robustness of behaviour in experimental ‘Beauty Contest’ games. *Economic Journal*, 107(445): 1684–1700
- Dworak, E.M., Revelle, W., Doebler, P., and Condon, D.M.** (2021). Using the International Cognitive Ability Resource as an open source tool to explore individual differences in cognitive ability. *Personality and Individual Differences*, 169: 109906
- Fe, E., Gill, D., and Prowse, V.** (2022). Cognitive skills, strategic sophistication, and life outcomes. *Journal of Political Economy*, 130(10): 2643–2704
- Ferraz, V., Pitz, T., Gardian, W., Kayar, D., and Sickmann, J.** (2025). Beyond Nash: Moral framing, level- k learning, and stationary concepts in experimental inspection games. *SSRN Working Paper 5373155*
- Fragiadakis, D.E., Kovaliukaite, A., and Rojo Arjona, D.** (2020). Does an individual predict opponents according to cognitive hierarchy? *Mimeo, University of Leicester*
- Frey, B.J. and Dueck, D.** (2007). Clustering by passing messages between data points. *Science*, 315(5814): 972–976
- Friedman, E. and Ward, J.** (2025). Stochastic choice and noisy beliefs in games. *Mimeo, Paris Schoof Economics*
- Fudenberg, D. and Liang, A.** (2019). Predicting and understanding initial play. *American Economic Review*, 109(12): 4112–4141
- García-Pola, B., Iriberri, N., and Kovářik, J.** (2020). Non-equilibrium play in centipede games. *Games and Economic Behavior*, 120: 391–433
- Georganas, S., Healy, P.J., and Weber, R.A.** (2015). On the persistence of strategic sophistication. *Journal of Economic Theory*, 159: 369–400
- Gill, D., Knepper, Z., Prowse, V., and Zhou, J.** (2025). How cognitive skills affect strategic behavior: Cognitive ability, fluid intelligence and judgment. *Games and Economic Behavior*, 149: 82–95
- Gill, D. and Prowse, V.** (2016). Cognitive ability, character skills and learning to play equilibrium: A level- k analysis. *Journal of Political Economy*, 124(6): 1619–1676
- Gill, D. and Prowse, V.** (2023). Strategic complexity and the value of thinking. *Economic Journal*, 133(650): 761–786
- Gill, D. and Rosokha, Y.** (2020). Beliefs, learning, and personality in the indefinitely repeated prisoner’s dilemma. *Centre for Advantage in the Global Economy (CAGE) Working Paper 489*
- Gill, D. and Rosokha, Y.** (2024). Beliefs, learning, and personality in the indefinitely repeated prisoner’s dilemma. *American Economic Journal: Microeconomics*, 16(3): 259–283

- Goeree, J.K.** and **Garcia-Pola, B.** (2023). S Equilibrium: A synthesis of (behavioral) game theory. *arXiv preprint 2307.06309v1*
- Goeree, J.K., Louis, P., and Zhang, J.** (2018). Noisy introspection in the 11–20 game. *Economic Journal*, 128(611): 1509–1530
- Gray, J. and Thompson, P.** (2004). Neurobiology of intelligence: Science and ethics. *Nature Reviews Neuroscience*, 5(6): 471–482
- Han, Y., Huffman, D., and Liu, Y.** (2026). Minds, models and markets: How managerial cognition affects pricing strategies. *Mimeo, Cornell University*
- Haruvy, E.** (2002). Identification and testing of modes in beliefs. *Journal of Mathematical Psychology*, 46(1): 88–109
- Hermes, H. and Schunk, D.** (2022). If you could read my mind – an experimental beauty-contest game with children. *Experimental Economics*, 25(1): 229–253
- Ho, T.H., Camerer, C., and Weigelt, K.** (1998). Iterated dominance and iterated best response in experimental “ p -beauty contests”. *American Economic Review*, 88(4): 947–969
- Ho, T.H. and Su, X.** (2013). A dynamic level- k model in sequential games. *Management Science*, 59(2): 452–469
- Ho, T.H., Park, S.E., and Su, X.** (2021). A Bayesian level- k model in n -person games. *Management Science*, 67(3): 1622–1638
- Hyndman, K., Özbay, E.Y., Schotter, A., and Ehrblatt, W.** (2012). Belief formation: an experiment with outside observers. *Experimental Economics*, 15(1): 176–203
- Hyndman, K., Terracol, A., and Vaksman, J.** (2022). Beliefs and (in)stability in normal-form games. *Experimental Economics*, 25(4): 1146–1172
- Jin, Y.** (2021). Does level- k behavior imply level- k thinking? *Experimental Economics*, 24(1): 330–353
- Kawagoe, T. and Takizawa, H.** (2012). Level- k analysis of experimental centipede games. *Journal of Economic Behavior & Organization*, 82(2-3): 548–566
- Lahav, Y.** (2015). Eliciting beliefs in beauty contest experiments. *Economics Letters*, 137: 45–49
- Lin, P.H. and Palfrey, T.R.** (2024). Cognitive hierarchies for games in extensive form. *Journal of Economic Theory*, 220: 105871
- Lindner, F. and Sutter, M.** (2013). Level- k reasoning and time pressure in the 11–20 money request game. *Economics Letters*, 120(3): 542–545
- Nagel, R.** (1995). Unraveling in guessing games: An experimental study. *American Economic Review*, 85(5): 1313–1326
- Nyarko, Y. and Schotter, A.** (2002). An experimental study of belief learning using elicited beliefs. *Econometrica*, 70(3): 971–1005

- Polonio, L.** and **Coricelli, G.** (2019). Testing the level of consistency between choices and beliefs in games using eye-tracking. *Games and Economic Behavior*, 113: 566–586
- Proto, E., Rustichini, A., and Sofianos, A.** (2019). Intelligence, personality and gains from cooperation in repeated interactions. *Journal of Political Economy*, 127(3): 1351–1390
- Proto, E., Rustichini, A., and Sofianos, A.** (2022). Intelligence, errors, and cooperation in repeated interactions. *Review of Economic Studies*, 89(5): 2723–2767
- Raven, J., Raven, J.C., and Court, J.H.** (2000). *Manual for Raven’s Progressive Matrices and Vocabulary Scales*. San Antonio: Pearson
- Rey-Biel, P.** (2009). Equilibrium play and best response to (stated) beliefs in normal form games. *Games and Economic Behavior*, 65(2): 572–585
- Romero, J. and Rosokha, Y.** (2018). Constructing strategies in the indefinitely repeated prisoner’s dilemma game. *European Economic Review*, 104: 185–219
- Romero, J. and Rosokha, Y.** (2023). Mixed strategies in the Indefinitely Repeated Prisoner’s Dilemma. *Econometrica*, 91(6): 2295–2331
- Schipper, B.C. and Zhou, H.** (2024). Level-k thinking in the extensive form. *Economic Theory*, 1–41
- Selten, R.** (1967). Die strategiemethode zur erforschung des eingeschränkt rationalen verhaltens im rahmen eines oligopolexperiments. In H. Sauermann, editor, *Beiträge zur Experimentellen Wirtschaftsforschung*, 136–168. J.C.B. Mohr, Tübingen
- Sofianos, A.** (2025). Intelligence in experimental economics. *Oxford Research Encyclopedia of Economics and Finance*, doi.org/10.1093/acrefore/9780190625979.013.853
- Spiliopoulos, L.** (2012). Pattern recognition and subjective belief learning in a repeated constant-sum game. *Games and Economic Behavior*, 75(2): 921–935
- Stahl, D.O.** (1996). Boundedly rational rule learning in a guessing game. *Games and Economic Behavior*, 16(2): 303–330
- Stahl, D.O. and Wilson, P.W.** (1994). Experimental evidence on players’ models of other players. *Journal of Economic Behavior & Organization*, 25(3): 309–327
- Wunder, M., Littman, M., and Stone, M.** (2009). Communication, credibility and negotiation using a cognitive hierarchy model. *Proceedings of the AAMAS 2009 Workshop on Multi-agent Sequential Decision-Making in Uncertain Domains*, 73–80

Appendix A:

Supplemental Appendix

(Intended for Online Publication)

Appendix A.1 Further figures and tables

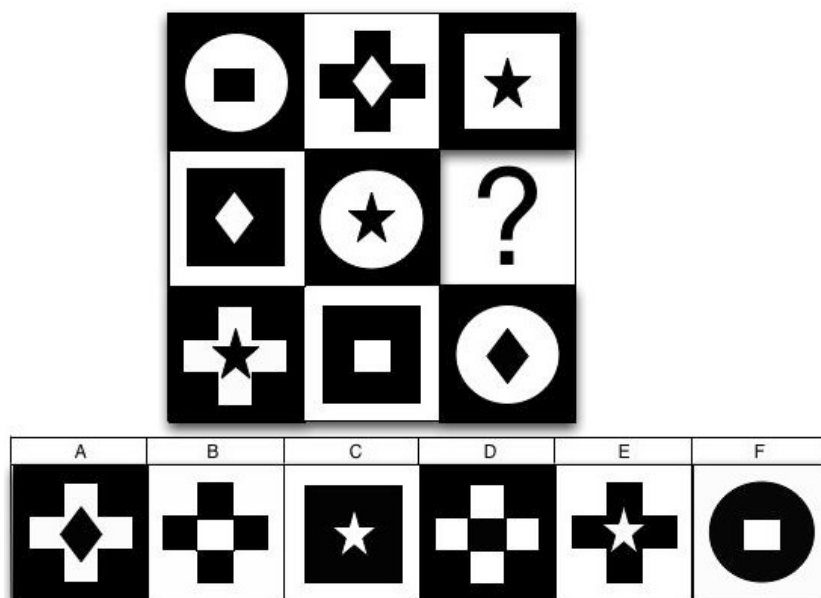
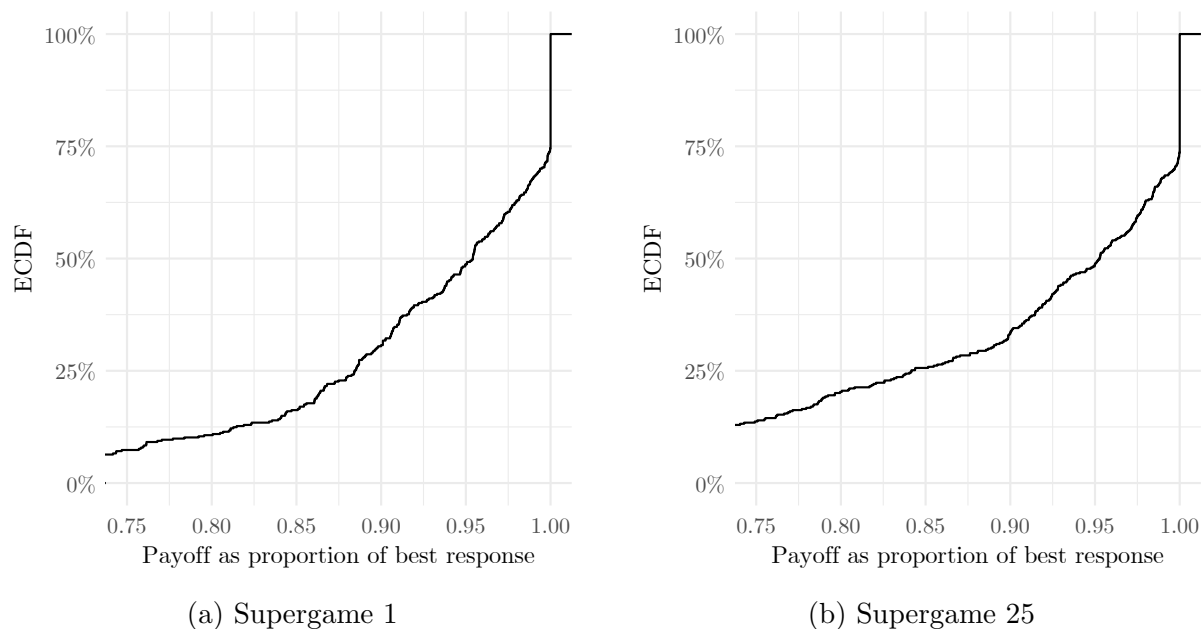


Figure A.1: Sample item from test of matrix reasoning



Notes: The figure shows the empirical cumulative distribution function (ECDF) across subjects of the expected payoff of the subject's chosen strategy, given her beliefs, as a proportion of the expected payoff of the best response to her beliefs. Before calculating expected payoffs, we normalize payoffs to range from 0 to 1.

Figure A.2: ECDF of payoffs as proportion of best response to beliefs

Cluster	N	<i>AD</i>	<i>DG</i>	<i>DTFT</i>	<i>RAND</i>	<i>G</i>	<i>2TFT</i>	<i>TFT</i>	<i>G2</i>	<i>TF2T</i>	<i>AC</i>
1	56	10	10	10	10	10	10	10	10	10	10
2	34	5	9	9	5	10	17	16	17	9	3
3	31	18	9	11	12	9	9	8	10	8	6
4	30	50	8	10	5	5	5	5	5	5	2
5	25	20	15	10	5	10	10	10	5	5	10
6	21	5	20	15	5	10	10	10	10	10	5
7	20	10	10	20	5	10	10	15	5	10	5
8	15	23	10	5	5	5	7	5	10	7	23
9	14	40	10	10	5	5	5	5	5	10	5
10	13	10	10	10	5	10	10	10	5	25	5
11	13	75	3	3	3	3	3	3	3	3	1
12	12	32	5	5	2	5	5	5	5	5	31
13	12	30	15	20	5	5	5	5	5	5	5
14	11	5	5	5	5	10	5	20	5	10	30
15	10	91	1	1	1	1	1	1	1	1	1
16	10	49	3	4	3	3	3	5	3	2	25
17	8	5	5	5	0	5	5	5	5	10	55
18	8	15	8	15	3	25	7	10	10	5	2
19	8	10	10	7	1	10	10	10	30	7	5
20	5	5	8	10	0	40	15	10	5	5	2
21	4	10	5	10	0	5	50	10	5	5	0
22	4	10	0	0	0	0	0	0	0	0	90
23	4	35	0	0	0	0	0	0	0	0	65
24	4	15	1	8	2	1	1	50	1	6	15
25	3	20	50	7	2	5	2	2	5	5	2
26	3	10	0	0	0	70	0	0	0	0	20
27	3	38	0	0	3	21	2	28	0	0	8
28	3	10	2	2	2	2	39	39	2	2	0
29	3	5	3	5	1	10	2	2	1	70	1
30	2	0	60	0	0	40	0	0	0	0	0
31	1	50	0	0	0	0	50	0	0	0	0
32	1	1	1	1	1	1	1	91	1	1	1
33	1	0	100	0	0	0	0	0	0	0	0
34	1	20	0	80	0	0	0	0	0	0	0
35	1	30	0	0	50	0	0	0	0	0	20

394

Table A.1: Belief clusters in Supergame 1 (all clusters)

Cluster	N	<i>AD</i>	<i>DG</i>	<i>DTFT</i>	<i>RAND</i>	<i>G</i>	<i>2TFT</i>	<i>TFT</i>	<i>G2</i>	<i>TF2T</i>	<i>AC</i>
1	44	20	15	15	5	10	10	10	5	5	5
2	36	10	10	10	10	10	10	10	10	10	10
3	33	5	10	10	5	15	15	15	10	10	5
4	32	50	5	10	5	5	5	5	5	5	5
5	27	2	15	20	2	11	11	11	11	15	2
6	17	25	25	20	3	5	5	5	5	5	2
7	17	80	5	2	1	1	2	5	2	2	0
8	15	98	0	0	0	0	1	1	0	0	0
9	15	15	5	20	5	5	5	30	5	5	5
10	14	10	5	5	2	20	30	15	5	5	3
11	14	10	30	40	1	1	9	8	1	0	0
12	14	5	20	15	5	10	10	20	5	5	5
13	11	40	40	3	2	3	3	3	2	3	1
14	11	65	5	20	0	0	3	5	2	0	0
15	11	25	3	25	3	30	3	3	2	3	3
16	10	40	3	3	3	3	3	2	3	3	37
17	9	5	40	10	5	5	5	10	5	5	10
18	8	0	30	10	0	10	25	10	15	0	0
19	8	40	5	40	0	4	2	5	2	2	0
20	7	1	5	40	7	5	3	30	3	5	1
21	6	0	5	0	0	10	10	10	5	10	50
22	6	25	55	2	0	5	5	8	0	0	0
23	5	30	0	0	0	0	0	70	0	0	0
24	5	30	10	20	0	0	30	10	0	0	0
25	3	0	0	0	0	100	0	0	0	0	0
26	3	20	1	1	1	1	1	30	1	1	43
27	3	0	15	5	0	60	0	20	0	0	0
28	3	30	0	0	0	70	0	0	0	0	0
29	2	0	10	0	0	0	0	0	0	0	90
30	2	0	5	0	0	0	5	25	50	10	5
31	1	0	0	0	0	0	0	50	0	50	0
32	1	0	0	100	0	0	0	0	0	0	0
33	1	1	1	1	1	1	1	1	1	91	1
394											

Table A.2: Belief clusters in Supergame 25 (all clusters)

Classification	AD	DG	DTFT	RAND	G	2TFT	TFT	G2	TF2T	AC
R = 32										
L1	14	15	15	6	11	12	11	6	7	4
L2	50	16	14	2	3	4	5	2	2	1
Neither	21	14	8	5	6	10	10	5	11	10
R = 40										
L1	12	12	13	7	11	11	11	8	9	6
L2	38	7	11	2	12	7	12	2	3	6
Neither	14	21	22	3	9	8	11	6	4	3
R = 48										
L1	9	10	11	4	13	13	11	11	11	8
L2	6	5	4	1	25	11	13	9	9	17
Neither	23	12	15	4	10	9	11	5	5	5

(a) Beliefs in Supergame 25 (%)

Classification	AD	DG	DTFT	RAND	G	2TFT	TFT	G2	TF2T	AC
R = 32										
L1	23	23	30	2	0	11	9	2	0	0
L2	61	21	11	0	0	2	5	2	0	0
Neither	8	8	4	12	8	29	4	8	8	8
R = 40										
L1	11	3	0	3	11	11	23	23	9	6
L2	38	7	5	0	16	13	11	4	4	2
Neither	25	28	33	3	3	3	6	0	0	0
R = 48										
L1	0	3	0	0	17	20	11	26	17	6
L2	0	0	0	0	28	12	12	20	8	20
Neither	28	16	20	12	3	8	9	0	0	3

(b) Strategy choices in Supergame 25 (%)

Notes: See the notes to Table 4: we use the same methodology, now for Supergame 25.

Table A.3: Empirical beliefs and strategies by treatment-category pair in Supergame 25

Appendix B:

Experimental Screenshots

(Intended for Online Publication)

Match #1 of 25

Plans	My Plan
#1 Choose randomly between A and B in every round. At the beginning of every round, the computer flips a computerized fair coin for you: when your coin comes up heads, you choose A; when your coin comes up tails you choose B.	Select
#2 Choose A in round 1. In round 2: choose A if the other chose A in round 1; choose B if the other chose B in round 1. After round 2: choose A if the other chose A in both of the previous two rounds; choose B if the other chose B in either of the previous two rounds.	Select
#3 Choose B in round 1. After round 1: choose A if the other chose A in the previous round; choose B if the other chose B in the previous round.	Select
#4 Choose B in every round.	Select
#5 Choose A in round 1. After round 1: choose A if the other chose A in the previous round; choose B if the other chose B in the previous round.	Select
#6 Choose A in round 1. After round 1: choose A if the other chose A in every one of the previous rounds; choose B if the other chose B in one or more of the previous rounds.	Select
#7 Choose A in rounds 1 and 2. After round 2: choose A if the other has never chosen B twice in a row (i.e., if the other has never chosen B in two consecutive previous rounds); choose B if the other has ever chosen B twice in a row.	Select
#8 Choose A in rounds 1 and 2. After round 2: choose A if the other chose A in either of the previous two rounds; choose B if the other chose B in both of the previous two rounds.	Select
#9 Choose A in every round.	Select
#10 Choose B in round 1. After round 1: choose A if the other chose A in every one of the previous rounds; choose B if the other chose B in one or more of the previous rounds.	Select

The table below describes the payoffs from the four pairs of choices that are possible in each round of a match.

My Choice in Round X	A	A	B	B
Other's Choice in Round X	A	B	A	B
My Payoff in Round X	32	12	50	25
Other's Payoff in Round X	32	50	12	25

Remember that at the end of each round the computer rolls a four-sided fair dice. The match ends when the computer rolls a 4.

	My Plan: Plan #4			
Dice Roll	4	1	3	2
Round	4	3	2	1
My Choice	B	B	B	B
Other's Choice	B	B	A	A
My Payoff	25	25	50	50
Other's Payoff	25	25	12	12

Match # 1 Summary

Match duration: **4** rounds

My total payoff: **150** points

Other's total payoff: **74** points

Next

Notes: The order that the ten strategies appeared on the subject's screen was randomized across subjects. In this example, the order of the strategies is: RAND; 2TFT; DTFT; AD; TFT; G; G2; TF2T; AC; DG.

Question Before We Implement Your Plan in Match # 1

Before we implement your plan in Match #1, **on the next screen we will ask you what you think the chances are that the other participant chooses each of the ten plans.**

In other words, we will ask you how many times out of 100 you think the other participant in Match #1 would choose each of the ten plans.

Of course, the other participant chooses his/her plan for Match #1 only once (just like you), not 100 times. But you can think of the question as a way of asking how likely the other participant is to choose each of the plans.

You will be paid according to the accuracy of your answer. You will make the most money on average if you answer truthfully what you think the chances are that the other participant chooses each of the plans.

You will not be paid for your answer until the end of the experiment. Your answer will not be shown to any other participant. Your answer will not affect the experiment in any way.

You do not need to understand the details of how the payment works. If you are interested, the details follow below. If you are not interested, you can stop reading this screen now.

Details of how the payment works:

For this task, you start with 400 points and you will then be penalized for inaccuracy. The total penalty can never be more than 400 points. Your payment for this task will be 400 points minus the total penalty.

For each of the ten possible plans, the penalty will be calculated as follows.

- 1. If that plan is chosen by the other participant, then the penalty will be smaller the higher your answer about the chances that the other participant chooses that plan. In that case, the penalty is calculated as follows:*

$$(100 - \text{Your Answer}) \times (100 - \text{Your Answer}) \times (0.02 \text{ points})$$

- 2. If that plan is not chosen by the other participant, then the penalty will be smaller the lower your answer about the chances that the other participant chooses that plan. In that case, the penalty is calculated as follows:*

$$(\text{Your Answer}) \times (\text{Your Answer}) \times (0.02 \text{ points})$$

Next

Question Before We Implement Your Plan in Match # 1

Plans	Other's Plan (Chance, %)
#1 Choose randomly between A and B in every round. At the beginning of every round, the computer flips a computerized fair coin for you: when your coin comes up heads, you choose A; when your coin comes up tails you choose B.	
#2 Choose A in round 1. In round 2: choose A if the other chose A in round 1; choose B if the other chose B in round 1. After round 2: choose A if the other chose A in both of the previous two rounds; choose B if the other chose B in either of the previous two rounds.	
#3 Choose B in round 1. After round 1: choose A if the other chose A in the previous round; choose B if the other chose B in the previous round.	
#4 Choose B in every round.	
#5 Choose A in round 1. After round 1: choose A if the other chose A in the previous round; choose B if the other chose B in the previous round.	
#6 Choose A in round 1. After round 1: choose A if the other chose A in every one of the previous rounds; choose B if the other chose B in one or more of the previous rounds.	
#7 Choose A in rounds 1 and 2. After round 2: choose A if the other has never chosen B twice in a row (i.e., if the other has never chosen B in two consecutive previous rounds); choose B if the other has ever chosen B twice in a row.	
#8 Choose A in rounds 1 and 2. After round 2: choose A if the other chose A in either of the previous two rounds; choose B if the other chose B in both of the previous two rounds.	
#9 Choose A in every round.	
#10 Choose B in round 1. After round 1: choose A if the other chose A in every one of the previous rounds; choose B if the other chose B in one or more of the previous rounds.	

The sum must add up to 100:

What do you think the chances are that the other participant chooses each of the ten plans?

In other words, how many times out of 100 do you think the other participant in Match #1 would choose each of the ten plans?

- Please use only whole numbers for your answer.
- You will be paid according to the accuracy of your answer. You will make the most money on average if you answer truthfully what you think the chances are that the other participant chooses each of the plans.
- You will not be paid for your answer until the end of the experiment.
- Your answer will not be shown to any other participant.
- Your answer will not affect the experiment in any way.