

Escuela de Negocios

Tipo de documento: Preprint



Collective problem decomposition improves the wisdom of deliberative crowds

Autorías: Barrera-Lemarchand, Federico; Lescano Charreau, Laura Victoria; Ruiz, Julieta; Cáceres, Nuria; Carrillo, Facundo; Sigman, Mariano; Navajas, Joaquín

Año: 2026

¿Cómo citar este documento?

Barrera-Lemarchand, F., Charreau, L. V. L., Ruiz, J., Cáceres, N., Carrillo, F., Sigman, M., & Navajas, J. (2026, March 28). Collective problem decomposition improves the wisdom of deliberative crowds. Retrieved from osf.io/preprints/psyarxiv/5xyze_v3

El presente Preprint se encuentra alojado en el Repositorio Digital Universidad Torcuato Di Tella para su preservación y archivo, bajo una licencia Creative Commons Atribución 4.0 según lo indicado en la fuente original del documento.

Dirección: <https://repositorio.utdt.edu/handle/20.500.13098/14268>

Collective problem decomposition improves the wisdom of deliberative crowds

Federico Barrera-Lermarchand^{1,2,3}, Victoria Lescano-Charreau^{1,2,4}, Julieta Ruiz¹, Nuria Cáceres¹, Facundo Carrillo^{3,5}, Mariano Sigman¹, & Joaquin Navajas^{1,2}

¹ Laboratorio de Neurociencia, Escuela de Negocios, Universidad Torcuato Di Tella, Buenos Aires, Argentina

² Comisión Nacional de Investigaciones Científicas y Técnicas (CONICET), Buenos Aires, Argentina

³ Facultad de Ciencias Exactas y Naturales, Universidad de Buenos Aires, Argentina

⁴ Laboratorio Interdisciplinario del Tiempo y la Experiencia (LITERA), Universidad de San Andrés, Buenos Aires, Argentina

⁵ Instituto de Investigación en Ciencias de la Computación (ICC), CONICET; Universidad de Buenos Aires; Argentina.

Corresponding author: Joaquin Navajas (joaquin.navajas@utdt.edu)

Abstract

Understanding when and why social interaction improves human judgment is a central question in the behavioural sciences. We examine whether collective accuracy improves when groups break complex estimation problems into simpler components and generate approximate intermediate estimates, a reasoning strategy we call Collective Fermi Estimation. Across three experiments analysing more than 1,000 online group deliberations in text chatrooms, we study the causal effects and linguistic signatures of this strategy. In Study 1 ($N = 500$), spontaneous use of problem decomposition, as rated by human annotators, predicts lower collective estimation error. Study 2 ($N = 240$) provides causal evidence: groups instructed to apply problem decomposition outperform groups instructed to combine their initial guesses. Study 3 ($N = 160$) shows that the benefits of the method are larger when applied collectively than individually. We also introduce a fully automated approach to detect Fermi-style reasoning in conversations based on large language models (LLMs). Using five state-of-the-art LLMs, we obtain ratings that strongly correlate with human annotations and predict collective accuracy across all studies. These findings identify problem decomposition as a key mechanism behind the wisdom of deliberative crowds and provide tools to detect and promote it.

1. Introduction

The aggregation of many lay estimates often outperforms individual expert judgments (De Condorcet, 1785; Galton, 1907). This phenomenon, popularly known as the “wisdom of crowds” (Surowiecki, 2005), has been applied to a wide range of problems, including predicting financial markets (Ray, 2006), forecasting geopolitical events (Mellers et al., 2014a), improving medical diagnoses (Kurvers et al., 2016), and even decoding the smell of molecules (Keller et al., 2017). Given the ubiquity of this effect and its broad applications across policymaking (Epp, 2017), medicine (Krockow et al., 2020), and finance (Lee et al., 2022), understanding why collective judgments sometimes enhance decision accuracy and at other times fail to do so has become a central challenge for the social and behavioural sciences (Prelec et al., 2017; Navajas et al., 2018; Kameda et al., 2022).

Previous research has examined the impact of social influence on the wisdom of crowds, yielding mixed results. Some studies suggest that it can lead to herding and reduce accuracy (Raafat et al., 2009; Lorenz et al., 2011; Madirolas & de Polavieja, 2015), whereas others show that social interaction can improve collective estimates (Gürçay et al., 2015; Mellers et al., 2014b; Bahrami et al., 2010; Juni & Eckstein, 2015). For instance, averaging the consensus estimates reached by small groups has been found to significantly outperform the wisdom of large, independent crowds (Navajas et al., 2018). Although this result has been replicated across multiple contexts (Dezecache et al., 2022; Espina Mairal et al., 2024; Pescetelli et al., 2021), the procedural mechanisms that determine when group consensus enhances or undermines collective accuracy remain poorly understood.

Based on theoretical considerations (Landemore & Page, 2015), groups can follow two distinct procedures to reach a collective estimate. One approach involves combining the initial private signals, for example, by taking the mean, the median, a confidence-weighted mean, or any other form of aggregation. A fundamentally different approach uses the deliberation process to develop a shared line of reasoning and generate a new estimate that is not merely a function of the initial individual judgments.

To illustrate these two strategies, consider a group of four individuals who independently estimate the height of the Eiffel Tower as 100, 200, 500, and 800 meters, and then come together to reach a shared judgment. The group might decide to discard the extreme values and average the two middle estimates, arriving at 350 meters, just 26 meters off the correct height. They could also average all four values, weight them by confidence, or use the median. What all these procedures have in common is that the final estimate is derived directly from the original individual judgments. Therefore, if groups follow this approach, they must exchange their initial estimates during deliberation.

Alternatively, the group might engage in a joint reasoning process to solve the problem, setting aside their initial estimates. For instance, they could discuss the structural features of the Eiffel Tower, reasoning that it resembles a 100-story building, with each floor about 3.5 meters high, again yielding an estimate of roughly 350 meters. Unlike the previous approach, it does not require sharing initial estimates. Because this strategy extends a well-established problem-solving technique for back-of-the-envelope calculations (attributed to Enrico Fermi) to a group setting (Anderson & Sherman, 2010; Weinstein & Adam, 2009; Nityananda, 2014), we refer to it as the *Collective Fermi Estimation* strategy.

Previous works have shown that consensus estimates produced by small groups can outperform the classical wisdom of crowds, suggesting that groups do more than merely

averaging their initial judgments (e.g., Navajas et al., 2018; Dezecache et al., 2022; Espina-Mairal et al., 2024). However, because the content of group discussions was not available in those works, it remained unclear whether groups relied on a mixture of aggregation rules or on a qualitatively different reasoning procedure. The present research was designed to identify the strategies that groups use to reach consensus and to determine which of these strategies are most effective for improving collective accuracy.

The Fermi estimation procedure

Fermi problems are estimation problems that require approximate reasoning in situations where exact information is unavailable. Although the “Fermi Method” is widely referenced across fields ranging from mathematics education (Albarracín, & Gorgorió, 2019) to psychology (Gomilsek et al., 2024), its definition is often described informally. Previous work consistently characterizes the method as a structured estimation procedure in which a complex quantity is approximated by decomposing it into simpler components that can be estimated separately (Albarracín, & Gorgorió, 2014; Albarracín, & Gorgorió, 2019; Abay & Filiz, 2020; Gomilsek et al., 2024; Albarracín & Ärlebäck, 2025).

To clarify this concept, we reviewed prior descriptions of the Fermi method in the literature (Table S1). Despite differences in terminology, these definitions converge on two core reasoning steps. First, the target quantity is decomposed into simpler subcomponents that can be estimated independently. Second, approximate values are proposed for those components using rough approximations or “guesstimates”, which are then recombined to produce a final estimate.

Based on this convergence in prior work, we adopt the following operational definition throughout this paper: “*The Fermi Method consists of breaking the problem into simpler components, making rough estimations for each component, and calculating the final answer based on those estimations.*” When this procedure is implemented through group deliberation, we call it Collective Fermi Estimation.

Building on prior evidence that collective reasoning can enhance performance beyond individual reasoning (Sperber & Mercier, 2012; Moshman & Geil, 1998; Trouche et al., 2014), we hypothesized that groups engaging in the Collective Fermi Estimation procedure would achieve lower error. To test this hypothesis, we developed a variety of methods to quantify the extent to which deliberation relies on this strategy, combining human ratings with natural language processing of group discussions.

Study 1 uses this measure to examine whether greater reliance on Collective Fermi Estimation predicts higher accuracy. Study 2 introduces an experimental manipulation, providing some groups with explicit instructions to apply problem decomposition and comparing their performance to that of groups instructed to aggregate their initial values. Study 3 compares individual and collective implementations of the same reasoning procedure, showing that the advantage emerges when the strategy operates at the collective level rather than in isolation.

2. Results

Study 1: Replication of Navajas et al. (2018) in chatrooms

500 participants (288 female, mean age 25.4 yr, s.d. 7.7 yr), organized into 125 groups of four individuals, participated in Study 1. Participants were incentivized to provide accurate answers (for details, see Methods). The study, which was exploratory in nature, consisted of three stages (Fig. 1A). During the first stage (stage i1), participants were asked individually a series of eight general-knowledge estimation questions (Table S2). In the second stage (stage c), they were divided into groups of four, and they were asked a random subset of four of the previous questions. They were instructed to deliberate these questions in a virtual chatroom, and asked to try to reach consensus. All conversations were recorded. In the last stage (stage i2), which followed exactly the same procedure as the first stage, participants could provide new estimates and confidence values.

Despite methodological differences with the original study, particularly with respect to group size (four vs. five) and modality (face-to-face vs. chatroom conversations), we replicated the observation that averaging collective estimates led to lower estimation error than averaging initial independent values (Figure 1B, see Methods for details). We observed this pattern for all values displayed in the x-axis of Figure 1B: For example, the estimation error of aggregating 8 randomly-chosen consensus estimates was substantially lower than the error obtained by averaging 32 randomly-chosen individual values (Mann–Whitney U test: $z=19.3$, $p=2 \times 10^{-83}$; effect size: Cohen's $d = 6.22$). This finding implies that collective estimates were more accurate than what we would have observed if groups had reached consensus by computing the mean of their initial values (Fig. 1C, normalized error of within-group averages: 1.18 ± 0.05 , normalized error of consensus estimates: 0.72 ± 0.04 , Wilcoxon signed-rank test: $z=7.3$, $p=3 \times 10^{-13}$; effect size: Cohen's $d = 0.95$).

Lastly, comparing participants' estimates from the first and third stages, separating questions that were deliberated in groups from those that were not, reveals a markedly greater reduction in error for the deliberated questions (Fig. 1D). Although participants showed a modest improvement in accuracy for undeliberated questions, the improvement was substantially larger for the deliberated questions (generalized linear mixed-effects model, see Methods for details: $\beta = -0.48$ [-0.53, -0.43], $SE = 0.03$, $t(2787) = -17.9$, $p = 5 \times 10^{-68}$, $R^2 = 0.148$).

In principle, one might assume that collective deliberation was dominated by individuals with higher accuracy or confidence, and that this alone could account for the improved accuracy of consensus estimates compared to simple averages. However, the data indicate that contributions within groups were approximately uniform. When we rank participants by individual error (Fig. S1A) or by confidence (Fig. S1B), the proportion of interventions made by each of the four group members is statistically indistinguishable from an equal share of 25% (two one-sided tests for equivalence, $p < 3 \times 10^{-5}$ for all comparisons; see Table S2 for details). This suggests that the accuracy gains were not simply driven by a few dominant individuals, but rather emerged from broadly distributed contributions within the group.

Collective Fermi Estimation is Associated with Lower Error

We then turned to the central research question of this study: What strategy do groups use to reach consensus? Do they share and combine their initial estimates, or do they recompute the value through problem decomposition? To investigate this, we trained ten individuals (8 female; mean age = 22.8 years, s.d. = 5.4) to rate each conversation along two dimensions using 11-point Likert scales (see Methods for rating instructions). The first rating assessed the extent to which group members shared and combined their initial judgments, a strategy we call “Sharing Numbers” (or simply “Numbers”). The second rating captured the extent to which participants generated new estimates through problem decomposition, a strategy we refer to as “Collective Fermi Estimation” (or simply “Fermi”).

Figure 2A illustrates two real examples: one conversation in which participants combined their initial estimates to reach consensus (left panel), and another in which participants decomposed the problem into intermediate quantities and proposed approximate values for those quantities (right panel). These examples illustrate how different conversational strategies can lead raters to assign different “Numbers” and “Fermi” scores. In particular, conversations dominated by the exchange and aggregation of initial estimates would typically receive higher “Numbers” ratings, whereas conversations involving problem decomposition and rough approximations would tend to receive higher “Fermi” ratings. These examples are intended only as illustrative cases showing how raters may interpret the conversational strategies present in the discussion, rather than implying that any specific numerical rating is uniquely correct for these conversations.

Study 1 (N=500)

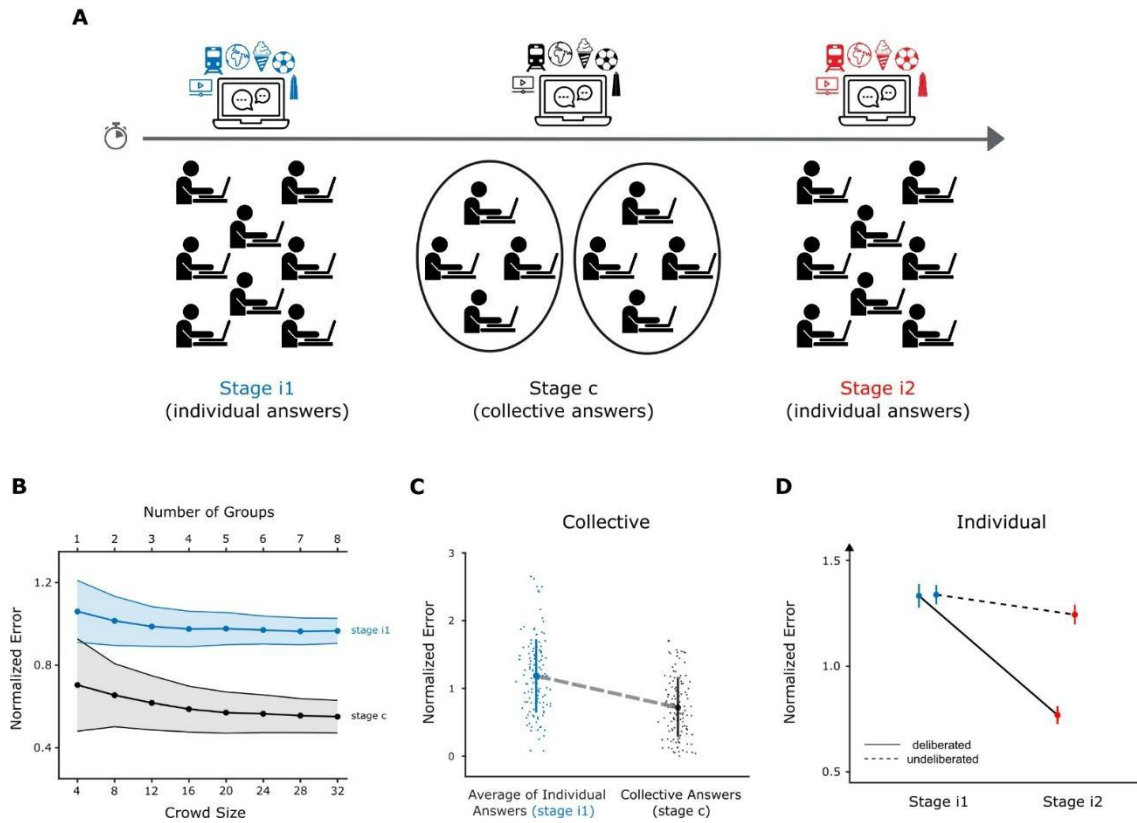


Fig. 1. Empirical results for Experiment 1. (A) The study had three stages. On Stage i1 participants answered eight general-knowledge questions individually. On Stage c, participants were divided into groups of four, and were asked half of the questions from the previous stage. Stage i3 followed the same procedure as Stage i1. (B) Normalized error of the average of n individual answers (blue line for stage i1), and normalized error of the average of $m = n/4$ collective estimates (black line, stage c). The error bars correspond to the standard deviation of the means. (C) Average normalized error and distributions of normalized errors of the individual answers (blue), and the collective answers (black). The error bars correspond to the standard deviation of the means. (D) Average normalized error for the individual stages, for both deliberated and undeliberated questions. The error bars correspond to 95% confidence intervals.

We observed moderate-to-strong pairwise Spearman correlations between raters, ranging from .35 to .76 (Fig. 2B; inter-rater correlation for “Numbers”: 0.518 ± 0.002 , $p < 6 \times 10^{-20}$; for “Fermi”: 0.661 ± 0.001 , $p < 3 \times 10^{-53}$), indicating a high level of consistency across them. Given this reliability, we used the average rating across raters as a summary measure of how strongly each strategy was employed in each conversation. These averaged ratings revealed a strong negative correlation between “Numbers” and “Fermi” scores (Fig. 2C; Pearson $r = -0.92$, $p = 6 \times 10^{-176}$), suggesting that the two strategies were typically used in a mutually exclusive manner.

We conducted two control analyses to increase confidence in human ratings. First, we found that the mean “Numbers” rating was positively correlated with the actual proportion of messages in which participants shared a number matching one of the initial individual estimates (Pearson $r = 0.32$, $p = 3.6 \times 10^{-11}$), and that this behavioral measure also correlated negatively with “Fermi” ratings (Pearson $r = -0.30$, $p = 3 \times 10^{-10}$).

Second, we conducted an additional annotation exercise designed to capture a broader form of deliberation. A separate group of six raters (5 female; mean age = 25.8 years, s.d. = 5.0) evaluated the conversations using more general instructions, rating the extent to which participants exchanged reasons to justify the group’s consensus value. This instruction was motivated by the distinction commonly made in the collective intelligence literature between number-based aggregation and deliberation involving the exchange of reasons (e.g., Landemore & Page, 2014). Because the Fermi method is itself a reasoning procedure, we expected that conversations exhibiting stronger Fermi-style reasoning would also receive higher reasoning ratings.

Consistent with this expectation, these broader reasoning ratings were strongly correlated with the original “Fermi” ratings (Pearson $r = 0.84$, $p = 2 \times 10^{-115}$), and the main results reported in Fig. 2 remain unchanged under this alternative annotation scheme (Fig.

S2). Importantly, these reasoning ratings were not intended as a mechanistic operationalization of the Fermi method itself, but to increase confidence in the robustness of the annotations: because Fermi-style estimation is one form of reason-based deliberation, similar results under this broader coding scheme indicate that the findings do not hinge on a narrowly framed annotation instruction.

We then examined whether and how the application of each strategy was associated with lower or higher estimation error. We found that conversations with higher “Numbers” ratings produced consensus estimates with higher estimation error (Fig. 2D, Pearson correlation: $r = 0.16$, $p = 0.004$) and that conversations scoring more highly on the “Fermi” strategy led to collective estimates which were less erroneous (Fig. 2E, Pearson correlation: $r = -0.20$, $p = 2 \times 10^{-4}$).

One might assume that groups with higher “Fermi” scores were more accurate at the collective stage simply because they were already more accurate at the individual stage. However, our data rule out this explanation. Conversations in which the Fermi score exceeded the Numbers score did not originate from more accurate individuals than those where the Numbers score was higher (Fig. 2F; Cohen’s $d = 0.018$; Bayes Factor $BF_{01} = 8.5$, strong evidence for the null hypothesis). The difference in accuracy between these groups only emerged after deliberation (Fig. 2F; Cohen’s $d = 0.37$; Bayes Factor $BF_{01} = 0.085$, strong evidence for the alternative hypothesis).

Another potential concern is that Fermi-based discussions may simply take longer than conversations focused on sharing numbers and this alone may drive performance. However, while conversations with higher Fermi ratings indeed tended to be longer (Pearson $r = 0.42$, $p = 6 \times 10^{-19}$), conversation length itself did not predict estimation error (Pearson $r = -0.07$, $p = 0.2$), suggesting that the observed performance gains are not simply driven by groups spending more time discussing the problem.

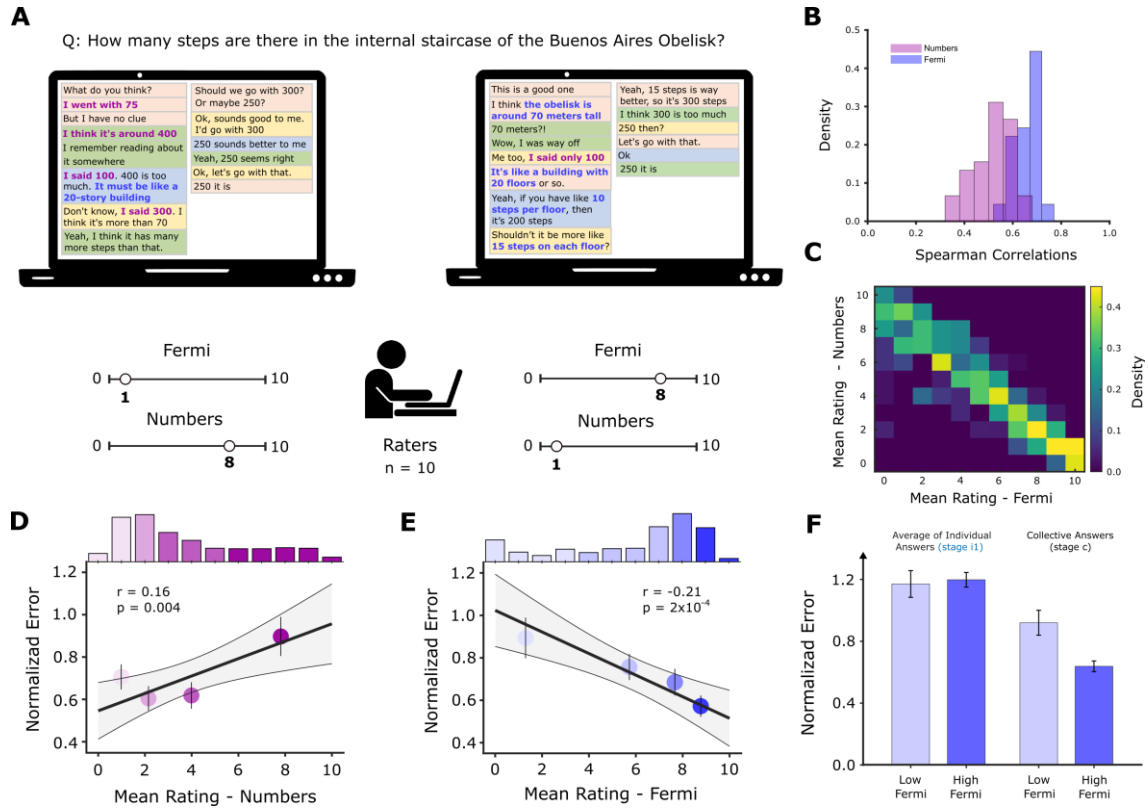


Fig. 2. Analysis of conversations. (A) Examples of conversations rated as high in either the “Numbers” strategy (sharing individual estimates, left panel) or the “Fermi” strategy (collective approximation, right panel). These examples are illustrative and do not pretend to imply that these particular values (1 and 8) are uniquely correct. Ten raters classified each conversation along these dimensions. (B) Distribution of pairwise Spearman correlations between raters for each strategy. (C) Density plot of mean “Numbers” and “Fermi” ratings, showing a strong negative correlation between strategies. (D–E) Normalized error as a function of average strategy ratings. Dots represent quartiles in the distribution of mean ratings, which is displayed above each plot. Lines depict best linear fits and shades represent 95% confidence intervals. (F) Normalized error before and after deliberation for conversations categorized as “low Fermi” (or “high Numbers”) or “high Fermi” (or “low Numbers”) based on which strategy received the higher average rating.

Predicting collective accuracy through automated language analysis

Because identifying groups that implemented the Collective Fermi Estimation strategy relies on human ratings, the process remains costly, time-consuming, and difficult to scale. To address this, we developed a computational method to automatically identify high-performing conversations. The objective of this analysis is not to provide a mechanistic test of Fermi reasoning, but rather to examine if a simple and low-cost linguistic proxy could identify conversational patterns associated with higher Fermi ratings and lower collective error. This motivation is relevant because extracting the wisdom of crowds from deliberative settings has numerous potential applications (e.g., prediction markets, financial forecasting, health decision-making, etc.) and scalable methods for identifying high-performing conversations could therefore be practically valuable.

The underlying intuition of this approach is simple: If groups use the deliberation stage to generate a new collective estimate, they are likely to use words that are semantically related to those in the estimation question. For example, consider a group tasked with estimating the number of steps in the internal staircase of the Buenos Aires Obelisk (Fig. 3A). A group decomposing the problem might compare the Obelisk to a 20-story building, estimate that a typical staircase has around 15 steps per floor, and conclude that the Obelisk should have approximately 250 steps, a solution which is reasonably close to the actual value of 206. Such a conversation would naturally include words that either appear in the question itself (e.g., “staircase,” “steps”) or are semantically related (e.g., “building,” “floor”).

To quantify this intuition, we used pre-trained word embeddings (Gutiérrez & Keith, 2019), specifically FastText embeddings trained in Spanish (Bojanowski et al., 2017). Word embeddings (vector representations of words) capture semantic relationships, such

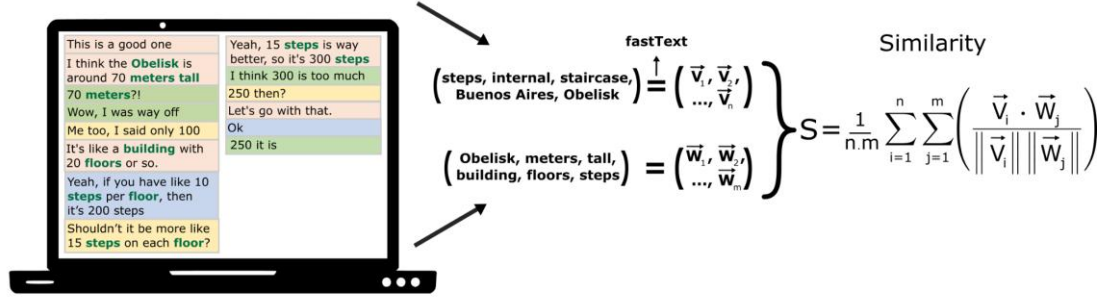
that semantically related words are represented by similar vectors. We computed the similarity between all words in the question and all words in the corresponding conversation, excluding stop words and short words. This yielded a distribution of similarity values, from which we extracted the mean to create a “Similarity” index (right panel in Fig. 3A).

We found that this Similarity index was positively correlated with “Fermi” ratings (Pearson $r = 0.21$, $p = 1 \times 10^{-5}$) and negatively correlated with “Numbers” ratings (Pearson $r = -0.16$, $p = 9 \times 10^{-4}$). Moreover, the Similarity index was negatively associated with estimation error (Fig. 3B; Pearson $r = -0.15$, $p = 0.006$), indicating that conversations more semantically related to the estimation problem tended to produce more accurate collective estimates. A mediation analysis further showed that this relationship was partially explained by Fermi ratings (Fig. 3C; total effect = -0.11 ± 0.03 , $p = 0.006$; direct effect = -0.08 ± 0.04 , $p = 0.02$; indirect effect = -0.03 ± 0.01 , $p = 0.009$).

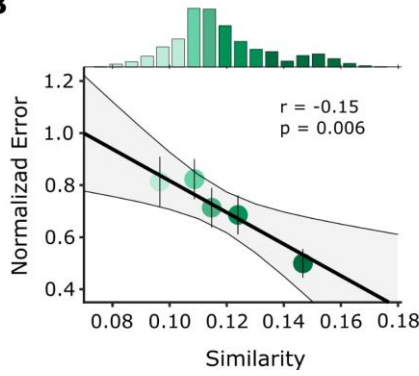
To assess whether simpler alternatives could reproduce the findings obtained with the Similarity index, we also examined two baseline methods: the total number of words in each conversation, and the number of words from the corresponding question (excluding stop words and short words) that appeared in the conversation. The former could reflect the complexity of the group’s reasoning process, while the latter offers a crude approximation of semantic overlap. However, neither measure was significantly associated with normalized group error (Pearson correlations: $r = -0.076$, $p = 0.17$; and $r = -0.082$, $p = 0.14$, respectively). These results suggest that simply counting words or matching lexical items from the question is insufficient to capture the underlying reasoning dynamics. Instead, the semantic similarity captured by word embeddings appears necessary to detect high-performing groups.

A

Q: How many **steps** are there in the **internal staircase** of the **Buenos Aires Obelisk**?



B



C

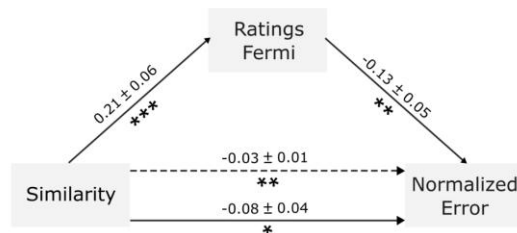


Fig. 3. Analysis of text. (A) For each conversation, we computed the similarity between all words in the conversation and all words in the question (excluding stop words and short words). Bags of words were created for both the question and the conversation, and FastText was used to obtain word embeddings. We then calculated the average similarity between all vectors from the two sets. (B) Normalized error as a function of similarity. Quintiles of the error distribution are shown, along with the best linear fit, 95% confidence intervals, and the Pearson correlation coefficient (r) with its corresponding p -value. Bars above the graph indicate the full distribution of similarity values. (C) Mediation analysis testing whether the relationship between similarity and normalized error is mediated by Fermi ratings.

The results presented so far provide evidence for an association between engaging in problem decomposition and producing collective estimates with higher estimation accuracy. To study if the implementation of this strategy, as opposed to just combining initial estimates, causally drives the observed reduction in collective error, we performed an experiment.

Study 2: Collective Fermi Estimation causally increases collective accuracy

240 participants (148 female, mean age 30.6 yr, s.d. 9.8 yr), organized in 60 groups of four individuals, participated in Study 2. Participants had monetary incentives to estimate these variables as accurately as possible (for details, see Methods). The protocol, hypotheses, and analysis plans were pre-registered at https://aspredicted.org/TL2_Q1Q.

The experiment (Fig. 4A) followed the same procedure as the previous study, with one key difference: after providing their initial estimates and before beginning the deliberation stage, participants watched a short instructional video on how to reach consensus (see Methods for full instructions). Half of the groups were randomly assigned to the “Fermi” condition, where they were instructed to break the problem into smaller estimation steps, make approximations, and recalculate the final answer. The other half were assigned to the “Numbers” condition, where they were instructed to share their initial estimates and combine them in any way they liked.

As predicted, we observed that groups in the “Fermi” condition produced collective estimates with lower error than those in the “Numbers” condition (Fig. 4B; Mann–Whitney U test: $z = 2.1$, $p = 0.04$; Cohen’s $d = 0.26$). To confirm that this difference reflected the implementation of different strategies, we trained ten naïve raters (4 female; mean age = 18.2 years, s.d. = 0.4), who had not participated in Study 1 and were blind to condition assignments, to evaluate the extent to which each conversation reflected the “Fermi” or “Numbers” strategy. As expected, conversations in the “Numbers” condition received higher “Numbers” ratings than those in the “Fermi” condition (Fig. 4C; Mann–Whitney U test: $z = 6.8$, $p = 1 \times 10^{-11}$; Cohen’s $d = 1.19$). Conversely, conversations in the “Fermi” condition received higher “Fermi” ratings than those in the “Numbers” condition (Fig. 4D; Mann–Whitney U test: $z = 6.3$, $p = 4 \times 10^{-10}$; Cohen’s $d = 1.06$).

Similarly, following our pre-registration, we found that instructing groups to share numbers led to a higher proportion of interventions containing numbers from the individual stage (Fig. 4E; Mann–Whitney U test: $z = 5.0$, $p = 6 \times 10^{-7}$; Cohen’s $d = 0.70$). Finally, the Similarity index (previously correlated with “Fermi” ratings in Study 1) was also higher for groups in the “Fermi” condition than in the “Numbers” condition (Fig. 4F; Mann–Whitney U test: $z = 2.5$, $p = 0.01$; Cohen’s $d = 0.34$). We also conducted an additional exploratory analysis to replicate Figs. 2D and 2E using the data of Study 2 (Fig. S3). We observed that, combining data from both conditions, collective error correlated negatively with Fermi ratings ($r = -0.24$, $p = .005$) and positively with Number ratings ($r = 0.22$, $p = .01$). Together, these results provide evidence that the instructions effectively modulated the strategies participants used to reach consensus.

We then conducted an additional exploratory analysis examining whether the collective estimate outperformed the individual estimates within each group. For each conversation, we ranked the four individual estimates according to their individual error (from most to least accurate) and computed the probability that the collective estimate was more accurate than each individual estimate (Fig. S4). Collective estimates always improved upon the least accurate individual estimate regardless of condition. However, differences between strategies emerged for more accurate individuals.

To quantify this pattern, we computed the fraction of initial individual estimates outperformed by the collective answer and compared this measure across conditions. A Wilcoxon rank-sum test indicated that collective estimates under the Fermi condition outperformed a larger fraction of initial estimates than those produced in the Numbers condition (Fermi: 0.73 ± 0.03 , Numbers: 0.63 ± 0.03 , $z = 2.33$, $p = 0.02$, Cohen’s $d = 0.44$). This suggests that the advantage of the Fermi condition lies not only in correcting poor

initial estimates, but also in more often surpassing the most accurate individual judgments in the group.

Discarding artefactual effects of instructions

To examine whether the observed advantage of the Fermi condition could be explained by mechanisms other than problem decomposition itself, we conducted additional exploratory analyses evaluating two alternative hypotheses. One possibility is that the Fermi instructions simply increased the diversity of information exchanged during discussion, for example by encouraging participants to articulate the reasons underlying their estimates. If this were the case, the observed gains in the Fermi condition might reflect greater informational diversity rather than problem decomposition per se.

To test this possibility, we implemented an information density metric previously proposed as a proxy for conversational informational diversity (Aceves & Evans, 2024). However, we observed that information density did not significantly predict collective error (Pearson correlation: $r = -0.08$, $p = 0.38$). These results suggest that the improved performance observed in the Fermi condition cannot be explained simply by increased informational diversity during discussion.

A second possibility is that the Fermi instructions improved performance by increasing coordination or semantic alignment among participants during deliberation. Following this hypothesis, we tested whether conversations in the Fermi condition exhibit stronger semantic convergence over time than those in the Numbers condition. Using the same FastText-based semantic similarity pipeline we previously presented (Fig. 3), we estimated the slope of message-to-message semantic similarity within each conversation. Contrary to the coordination hypothesis, we do not observe increasing semantic convergence over time in the Fermi condition (distribution of slopes, mean $\pm$ S.D. = -0.005 ± 0.03 , t-test: $t(89) = -1.5$, $p = 0.1$) and the best-fitting slopes do not predict

collective error (Pearson $r = 0.02$, $p = 0.82$). These results remain qualitatively unchanged when applying a 5-message sliding-window approach to compute local convergence patterns (distribution of slopes, mean $\pm$ S.D. = -0.0007 ± 0.009 , t-test: $t(89) = -0.8$, $p = 0.4$; correlation with collective error: Pearson $r = 0.07$, $p = 0.45$). Overall, these findings do not support the hypothesis that the performance gains observed in the Fermi condition are driven by semantic alignment between individuals during deliberation.

Linguistic signatures of Fermi ratings

To examine whether conversations identified as “Fermi” indeed reflect the reasoning processes described in its definition, we developed linguistic indicators capturing the two components of the method: problem decomposition and guesstimation. Problem decomposition was measured through expressions that break the target quantity into component structures (e.g., multiplicative or per-unit reasoning such as “per”, “each”, “every”, or “times”). Guesstimation was measured through linguistic markers reflecting approximate numerical reasoning under uncertainty, including expressions proposing numerical ranges (e.g., “between N and M”, where N and M are numbers) and approximate quantities (e.g., “approximately N”, “around N”, or “about N”, where N is a number).

Study 2 (N=240)

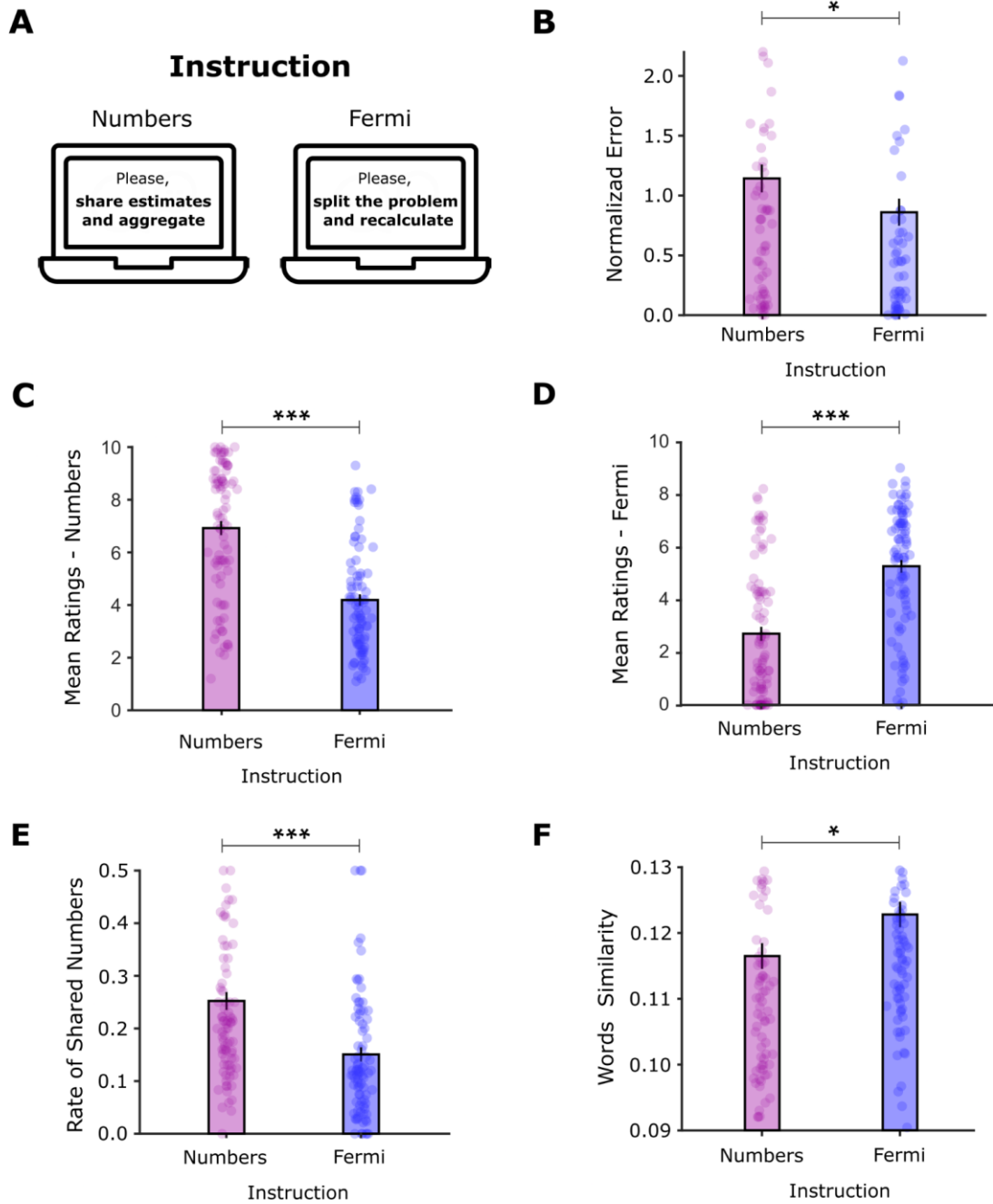


Fig. 4. Results of Experiment 2. (A) The experiment had the same procedure as Study 1, except that before the collective deliberation stage, participants received an instruction either to share numbers or to perform approximations and recompute the value. (B) Normalized error by instruction condition. (C) “Numbers” ratings provided by human raters by condition. (D) “Fermi” ratings provided by human raters by condition. (E) Proportion of shared numbers from the individual stage present in discussions, by condition. (F) Similarity index by instruction condition. In all cases, bars show the mean value, error bars depict s.e.m. and dots show data from individuals. Asterisks denote statistical significance (*: $p < .05$, **: $p < .01$, ***: $p < .001$).

Across studies, these linguistic indicators strongly predict the extent to which conversations exhibit Fermi-style reasoning. In Study 1, both decomposition and guesstimation markers significantly predict human-annotated Fermi ratings (Fig. 5A, $\beta_{PD} = 0.54$, $p = 5 \times 10^{-19}$; $\beta_{GE} = 0.43$, $p = 0.01$). In Study 2, the effect of the experimental manipulation (Fermi vs. Numbers instructions) on Fermi ratings is mediated by the increased presence of these markers in the conversations (Fig. 5B, total effect: 2.56 ± 0.36 , $p = 0.0006$, direct effect: 1.86 ± 0.37 , $p = 0.0005$, indirect effect of PD: 0.46 ± 0.12 , $p = 0.0004$; indirect effect of GE: 0.25 ± 0.11 , $p = 0.01$).

Taken together, these results indicate that conversations identified as exhibiting Fermi-style reasoning systematically contain the linguistic signatures predicted by the formal definition of the method. This provides converging evidence that the strategy implemented in the experiments reflects the combination of problem decomposition and rough approximation that characterizes the Fermi method.

While Study 2 offers new insights into the drivers of collective accuracy and supports the idea that groups using the Fermi method achieve better estimates, another explanation is possible. The observed reduction in error in the “Fermi” condition may not be due to collective problem decomposition itself, but rather to individual participants applying the Fermi method within the group, without necessarily engaging in a group-level reasoning process (Landemore & Mercier, 2010). In other words, the strategy might improve accuracy simply by being used, whether individually or collectively. To test this possibility, we conducted another experiment designed to disentangle individual and collective contributions to the effectiveness of the Fermi method.

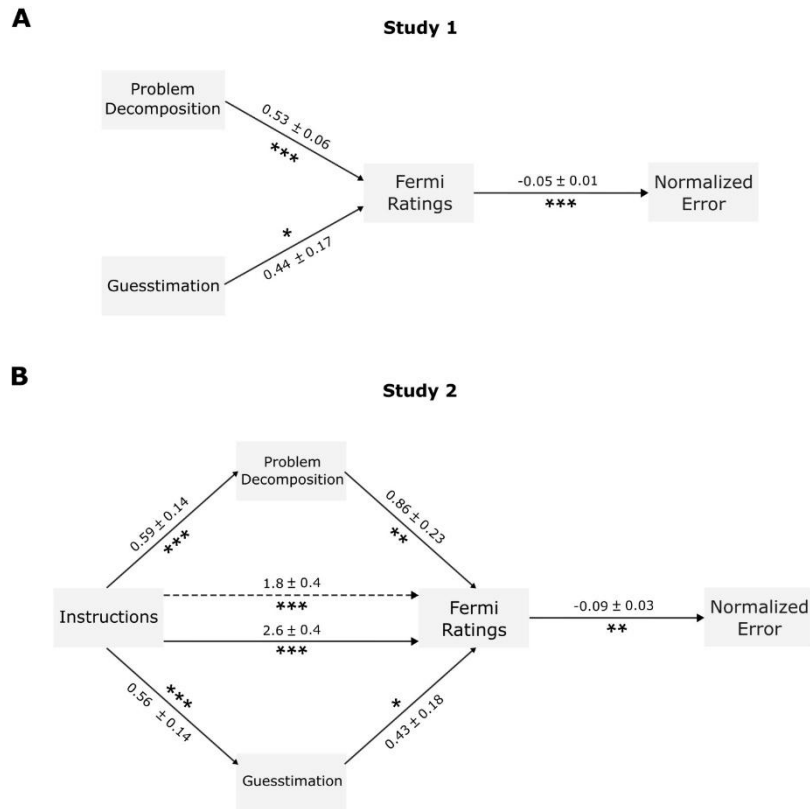


Fig. 5. Linguistic mechanisms underlying Fermi ratings. (A) Linguistic markers of problem decomposition and guesstimation in Study 1, derived from the formal definition of the Fermi method, predict human-annotated Fermi ratings. In turn, higher Fermi ratings are associated with lower normalized estimation error. Coefficients correspond to regression estimates ($\pm$ s.e.m.), and asterisks indicate statistical significance (*: $p < .05$, **: $p < .01$, ***: $p < .001$). (B) Mediation analysis for Study 2. Fermi instructions increase the presence of both decomposition and guesstimation markers in the conversations, which in turn predict higher Fermi ratings. Fermi ratings are negatively associated with normalized collective error. The dashed arrow indicates the direct effect of the experimental manipulation on Fermi ratings after accounting for the linguistic mediators.

Study 3: Fermi Estimation is more accurate in groups than alone

N=160 participants (95 female, mean age 29.8 yr, s.d. 10.9 yr) participated in Study 3 (pre-registered at https://aspredicted.org/SZ5_MNP). A random half of the participants were divided into 20 groups of 4 individuals, and the remaining 80 performed the study individually. For half of the participants, the experiment followed exactly the same procedure as the “Fermi” condition in the previous study: Participants received instructions to apply the Fermi method and deliberated in groups to reach a collective estimate (Fig. 6A). The key difference in this experiment was the addition of an individual condition, which applied to the other half of participants, who worked alone rather than in groups. These participants were instructed to apply the Fermi method individually and to write their reasoning and calculations in a chat window, without interacting with others (Fig. 6B). All written responses across both conditions were recorded (see Methods for details).

This design allowed us to directly compare the effectiveness of applying the Fermi method individually versus in groups as a strategy for improving accuracy (Figs. 6C and 6D). Following our pre-registration, we first tested whether final individual error at stage i2 differed between the Collective and Individual conditions. A mixed-effects linear regression with normalized individual error as the dependent variable, a dummy variable coding for the Collective condition, and random intercepts for the eight questions showed that participants in the Collective condition produced significantly lower errors at stage i2 than those in the Individual condition ($\beta = -0.24 \pm 0.07$, $p = 0.0007$). A Wilcoxon rank-sum test comparing the distributions of individual errors across conditions yielded the same conclusion ($z = 3.12$, $p = 0.002$, Cohen’s $d = 0.26$).

To facilitate interpretation of how deliberation affected performance relative to each participant’s initial estimate, we also report an exploratory analysis examining the

reduction in error between the initial individual estimates (stage i1) and the final individual estimates (stage i2). First, we focused on the collective condition (Fig. 6C). We observed that, while participants showed a modest improvement for non-deliberated questions (Cohen's $d = 0.08$), the improvement was substantially greater for collectively deliberated questions (Cohen's $d = 0.69$, $\beta = -0.45$ [$-0.59, -0.31$], $SE = 0.07$, $t(469) = -6.36$, $p = 5 \times 10^{-10}$, $R^2 = 0.10$).

We then examined the individual deliberation condition (Fig. 6D), where participants worked alone but were instructed to apply the Fermi method. As in the collective condition, participants showed a modest improvement for non-deliberated questions (Cohen's $d = 0.06$) and a larger improvement for the ones that underwent individual deliberation (Cohen's $d = 0.28$, $\beta = -0.17$ [$-0.28, -0.07$], $SE = 0.05$, $t(469) = -3.28$, $p = 0.001$, $R^2 = 0.02$). This suggests that applying the Fermi method individually may result in greater accuracy.

To formally compare conditions, we examined the reduction in error in the collective and individual conditions separately for deliberated and undeliberated items. For deliberated items, error reduction was substantially larger in the collective condition than in the individual condition (Mann–Whitney U test: $z = 6.3$, $p = 3 \times 10^{-10}$; Cohen's $d = 0.58$). No such difference was observed for undeliberated items (Mann–Whitney U test: $z = 1.0$, $p = 0.31$; Cohen's $d = 0.05$). These results suggest that, although the Fermi method improves accuracy when applied individually, collective deliberation provides an additional benefit beyond individual reasoning alone.

Study 3 (N=160)

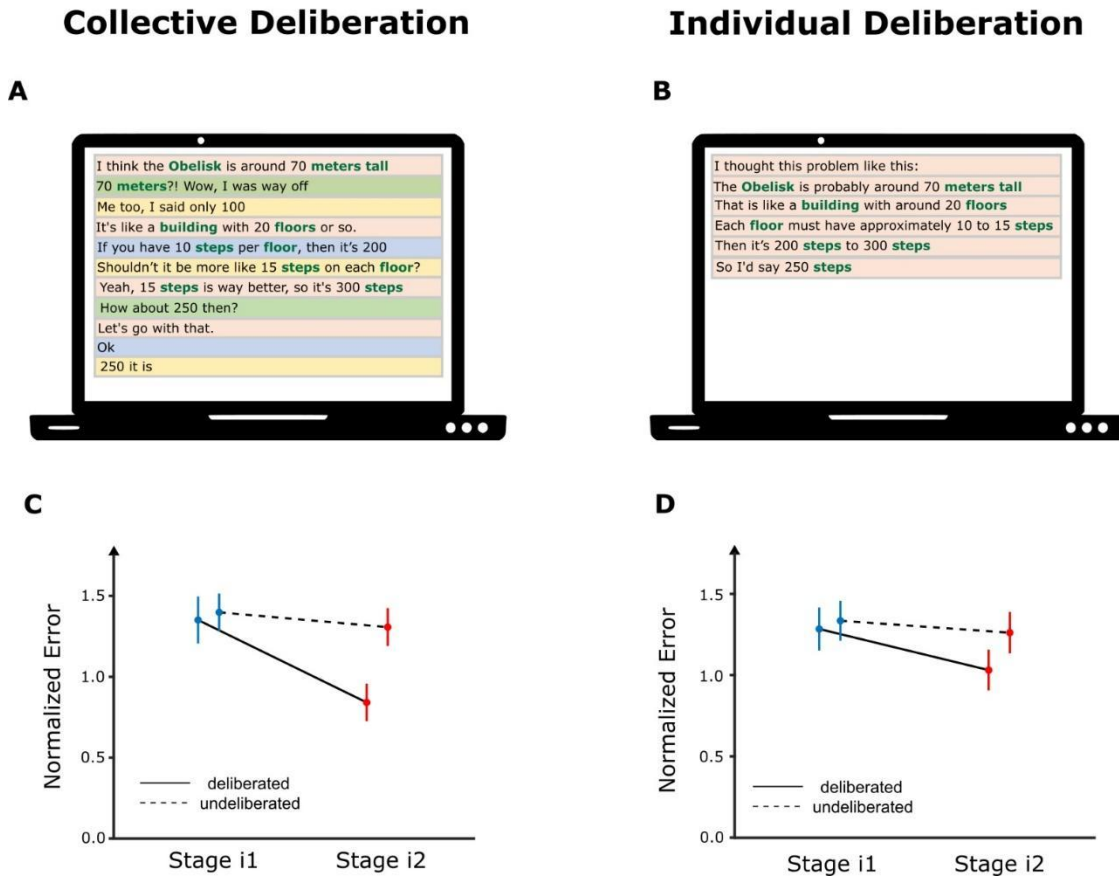


Fig. 6. Empirical results for Experiment 3. The procedure was the same as that of the previous experiment, with two exceptions: **(A)** half of the participants were divided into groups and received the original Fermi instruction, applying the Fermi method as a group. **(B)** the remaining half proceeded individually, and received a modified Fermi instruction prompting them to also apply the Fermi method individually (writing down their procedure). **(C)** Average normalized error for the individual stages, for both deliberated and undeliberated questions, for participants that in the collective condition. The error bars correspond to the 95% confidence intervals. **(D)** Average normalized error for the individual stages for participants that deliberated individually. We followed the same procedure as for the previous panel.

Why does the Fermi method benefit groups more than individuals?

Study 3 allowed us to examine whether the benefits of the Fermi method differ depending on whether the strategy is applied individually or collectively. In this experiment, all participants were instructed to apply the Fermi method, but some deliberated individually while others deliberated in groups. In both cases, participants improved their estimates relative to their initial responses but the magnitude of this improvement was substantially larger in the collective condition (Figs. 6C and 6D).

One possible explanation to this effect is that collective deliberation enabled metacognitive processes that are difficult to reproduce in isolation. Prior work in collective cognition suggests that interactive reasoning can benefit from supra-personal monitoring processes, such as detecting and correcting errors in others' reasoning (e.g., Bahrami et al., 2010; Frith, 2012; Shea et al., 2014). In the context of Fermi estimation, this mechanism is plausible because decomposing a problem generates multiple intermediate assumptions that can be re-evaluated during deliberation.

To test this possibility, we conducted an exploratory analysis of the language produced by participants in Study 3. We constructed a linguistic indicator capturing error-correction dynamics, based on expressions typically associated with evaluating or adjusting magnitudes (e.g., phrases like “too little” or “too much”, or evaluative corrections such as “better”). These expressions occurred substantially more frequently in collective deliberation than in individual reasoning (Fig. S5, unpaired t-test: $t(360)=-6.7$, $p=8 \times 10^{-11}$, Cohen's $d = 0.91$), and higher levels of such language were associated with lower estimation error in the subsequent individual stage (Pearson $r = -0.18$, $p = 0.0001$). Importantly, these markers did not differ between the Fermi and Numbers conditions in Study 2 (Fig. S5, unpaired t-test: unpaired t-test: $t(177)=0.9$, $p=0.4$, Cohen's $d = 0.14$), and did not predict collective error in Study 1 (Pearson $r = 0.31$, $p = 0.58$). This pattern

suggests that error-correction dynamics help explain why collective deliberation amplifies the benefits of the Fermi method, rather than why the Fermi method itself improves performance relative to number sharing.

Can large language models detect Fermi reasoning?

Because human annotation procedures can, in principle, be influenced by the instructions provided to raters, we conducted an additional analysis to evaluate whether Fermi ratings could be reproduced using an independent procedure based solely on the formal definition of the method (presented in the Introduction). Specifically, we generated AI-based Fermi ratings using large language models (LLMs) (Bermejo et al., 2025). For each conversation, the model was provided with the theoretical definition of the Fermi method and asked to rate, on a 0–10 scale, the extent to which the procedure in the conversation matched that definition.

To ensure that the results were not dependent on the idiosyncrasies of a particular model, we repeated this procedure using five different LLMs (gpt-4.1-mini, gpt-5-mini, gpt-5.2, claude-haiku-4.5, and claude-sonnet-4.6). Across models, the resulting ratings show strong correlations with the mean human annotations (Pearson $r > 0.7$, $p < 7 \times 10^{-64}$ for all comparisons, Table S3) as well as strong correlations among models themselves (Pearson $r > 0.8$, $p < 8 \times 10^{-98}$ for all comparisons, Table S3), indicating a high level of agreement across independent rating procedures.

Importantly, the main empirical relationships reported in this paper replicate when using AI-based ratings in place of the human annotations. Across all three studies, higher Fermi ratings generated by the language models are associated with lower estimation error (Table 1). In addition, these ratings show the same expected relationships with the

linguistic indicators introduced above: they are positively associated with both problem decomposition and guesstimation markers (Table 1).

These results provide converging evidence that the Fermi ratings capture a stable and identifiable reasoning pattern rather than an artefact of the human annotation procedure. Conversations rated as exhibiting stronger Fermi-style reasoning (whether by human raters or LLMs) systematically display the linguistic signatures predicted by the formal definition of the method and are associated with improved collective accuracy.

	Study 1			Study 2			Study 3		
	Error	PD	GE	Error	PD	GE	Error	PD	GE
GPT 4.1 mini	-0.20***	0.38***	0.13**	-0.21*	0.39***	0.33***	-0.24***	0.48***	0.13*
GPT 5.1 mini	-0.22***	0.39***	0.19***	-0.28**	0.44***	0.32***	-0.30***	0.45***	0.21***
GPT 5.2	-0.20***	0.45***	0.12*	-0.27**	0.45***	0.30***	-0.27***	0.50***	0.11*
Claude Haiku 4.5	-0.25***	0.35***	0.12*	-0.31***	0.40***	0.35**	-0.31***	0.35***	0.15**
Claude Sonnet 4.6	-0.22***	0.39***	0.15**	-0.29**	0.46***	0.34***	-0.32***	0.30***	0.05

Table 1. Replicating main results using LLM-based Fermi ratings. For each of five large language models (rows), the table reports Pearson correlation coefficients between LLM-generated Fermi ratings and key variables across the three studies. Columns show correlations with normalized estimation error (Error), problem decomposition markers (PD), and guesstimation markers (GE). Asterisks denote statistical significance (*: $p < .05$, **: $p < .01$, ***: $p < .001$).

Discussion

Across three studies, we provide converging evidence that small groups can reach more accurate estimates than large, independent crowds when they engage in problem decomposition, a strategy we term “Collective Fermi Estimation”. This approach involves making rough calculations and computing an approximate solution, rather than merely aggregating initial signals. In Study 1, we found that groups exhibiting higher use of this strategy, as determined by trained human raters, produced significantly more accurate estimates. Study 2 established a causal link: When groups were explicitly instructed to use the Collective Fermi Estimation strategy, their performance improved relative to groups prompted to aggregate their individual estimates. Finally, in Study 3, we demonstrated that the advantage of this strategy lies in its collective application: Aggregating the estimates of individuals who each applied the strategy independently yielded poorer performance than when the reasoning occurred through group deliberation (Landemore & Mercier, 2010).

Another contribution of this work is the development of a variety of computational methods to detect whether a group employed the Collective Fermi strategy based on the content of its conversation. Using word embeddings to measure the semantic similarity between words in the deliberation and those in the original estimation question, we constructed a low-cost index that correlates with human Fermi ratings and, critically, predicts group accuracy. This method outperformed simpler proxies such as word count or lexical overlap, suggesting that capturing the semantic structure of a group’s reasoning is essential for detecting meaningful deliberation. We also developed a pipeline to detect Fermi-style reasoning using LLMs and show that our results replicate using this fully automated approach. These results have practical implications as they provide new methods to extract information from crowdsourced estimates under social influence.

These findings also contribute to the broader literature on collective intelligence and the wisdom of crowds by helping to clarify when and how social interaction can improve collective accuracy. Prior work has shown that deliberation can either impair or enhance group performance, depending on the context and the nature of the interaction (Lorenz et al., 2011; Bahrami et al., 2010; Becker et al., 2017; Navajas et al., 2018). Here, we provide evidence in support of the hypothesis that deliberation grounded in reasoning (rather than information exchange) is a key mechanism driving superior group performance.

However, several limitations remain. As with much of the crowd wisdom literature, it is unclear whether our findings extend to domains where correct answers are categorical rather than continuous (for a notable exception, see Becker et al., (2022)). Moreover, many real-world group decisions involve not only categorical judgments, but also explanations, advice, planning, or the analysis of complex situations, for which numerical aggregation procedures do not directly apply. Our findings therefore speak most directly to a relatively narrow class of estimation problems. However, the mechanism we study (i.e., collective decomposition of a complex problem into simpler components) may generalize beyond numerical estimation. In principle, many complex judgments can be structured as a set of simpler sub-questions whose answers can then be integrated into a final decision. Future work should examine whether collective Fermi-style reasoning can provide similar benefits in domains where the outcome is not a numerical estimate but rather a categorical judgment, explanation, recommendation, or plan of action.

Another limitation of the present design is that, in Studies 1 and 2, questions that were deliberated differed from those that were not deliberated not only because they involved group interaction, but also because participants could devote more attention to them. Because deliberated questions were discussed collectively for up to five minutes, whereas undeliberated questions were answered individually within the overall task time, the two

types of questions may differ in the amount of individual attention allocated to them. As a result, comparisons between deliberated and undeliberated questions in those studies (which we present only in Fig. 1D) cannot cleanly isolate the effect of group interaction from differences in attention per question. Importantly, however, the central comparisons of the present work do not rely on this contrast. In Study 1, the key analyses focus on variation in how groups reason about deliberated questions. In Study 2, the main comparison is between two instruction conditions applied to deliberating groups. Finally, Study 3 provides a cleaner test of the role of social interaction by comparing individuals deliberating alone with groups deliberating together, holding the opportunity for deliberation constant. Together, these analyses suggest that while differences in attention may contribute to some contrasts, the central findings regarding collective Fermi-style reasoning are not driven by this feature of the design.

Beyond hypothesis testing, we believe our findings carry broader implications for the design of deliberative environments. The wisdom of crowds has long been interpreted as a model for the epistemic benefits of democratic judgment (Galton, 1907; Landemore, 2012; Surowiecki, 2005). Within this framework, a central theoretical distinction has been drawn between “voting mechanisms,” which aggregate individual preferences, and “deliberative mechanisms,” which emphasize shared reasoning and mutual justification (Landemore & Mercier, 2010). While deliberative democracy rests on the premise that genuine deliberation can improve collective decisions over information sharing (Landemore, 2017), empirical demonstrations of this principle have remained limited. Our findings help fill this gap by showing not only that groups naturally engage in collective reasoning under certain conditions, but also that they can be prompted to adopt such a mindset, yielding measurable epistemic benefits.

Overall, this work shows that collective reasoning, and in particular reasoning by approximation, underlies enhanced collective accuracy during deliberation, and offers new tools to detect and foster such processes, contributing to basic understanding of human group decision-making and informing practical applications aimed at improving crowdsourcing strategies.

Methods

Participants

Participants were recruited from the Participant Database at the Laboratory of Neuroscience of Universidad Torcuato Di Tella. They were informed that their participation was completely voluntary, and that they could withdraw their participation at any time. All data were completely anonymous. Overall, we collected data from 900 individuals, all residents of Argentina, across three samples: Study 1 involved 500 participants (288 female, mean age 25.4 yr, s.d. 7.7 yr), in Study 2 we recruited 240 participants (148 female, mean age 30.6 yr, s.d. 9.8 yr), and Study 3 involved 160 participants (95 female, mean age 29.8 yr, s.d. 10.9 yr).

In all cases, participants were entered in a draw for 5 prizes of 100 USD each. This procedure led to a participation fee with an expected value of roughly 6 USD per hour of experiment, which is a fair rate based on the guidelines of subject participation in Buenos Aires, Argentina. All protocols were approved by the ethics committee at Universidad Abierta Interamericana (Buenos Aires, Argentina), protocol 0-1108.

Study 1

The study consisted in three stages. In the first stage, participants independently responded to a series of eight numerical estimation questions within a 5-minute time limit. Alongside their answers, they indicated their level of confidence using a slider bar ranging from low (0) to high (100). Moving on to the second stage, participants communicated

via chat to try to reach a consensus on values for four of the eight questions from the first stage. Participants were instructed to try to reach consensus with a 5-minute time limit. If consensus was achieved earlier, participants had to wait until the 5 minutes elapsed before moving on to the next question. During this stage, participants were not asked to report their confidence levels. The third stage of the study mirrored the first stage, wherein participants individually tackled the same numerical estimation questions within a time limit of five minutes.

Study 2

Study 2 was identical to Study 1, except that they were given instructions on how to reach consensus. In addition to the previous procedure, participants were allocated to one of two conditions. In the *Numbers Condition*, they were explicitly instructed to “share their initial numerical estimates”. In the *Fermi Condition* they were explicitly instructed to “divide the problem into smaller estimation problems and perform all relevant computations to reach their final answer”. All hypotheses and data analysis plans were preregistered at https://aspredicted.org/TL2_Q1Q.

Study 3

In Study 3, all participants received the same instructions as in the Fermi Condition from the previous study. However, a random half were assigned to groups of four, while the other half worked individually. Those in the individual condition were asked to apply the Fermi Method on their own, documenting their procedure in the same chatroom used by the groups, though they remained alone for the entire experiment. This study was preregistered at https://aspredicted.org/SZ5_MNP.

Estimation questions

The eight questions were presented in random order in the first stage. A random half of them were selected for stage 2. For the third stage, the questions were presented in the

same order as the first stage. The complete set of questions and correct answers are available at Supplementary Table 1.

Supervised metrics

For this approach, we trained ten raters (8 female; mean age = 22.8 years, s.d. = 5.4), who were naïve to the hypotheses of the study, on identifying conversations that used the “Fermi” and “Numbers” strategies. Then, they individually read each conversation and determined how much of the Fermi method was actually used on a scale from 0 (completely absent) to 10 (completely present). They were also asked to rate how much the groups implemented the “Numbers” strategy, consisting on exchanging their answers and combining their initial numerical estimates to obtain a final group answer. The order of presentation of the explanation of the two strategies and the conversations to be rated were randomized across raters. The detailed instructions provided to the raters, including the alternative definition of the “Fermi” method, are available in Supplementary Note 1.

Similarity measures

We computed the semantic similarity between the conversation and the task using natural language processing (Manning, & Schutze, 1999; Carrillo et al., 2015). To this aim, we leverage pre-trained word embeddings, specifically FastText trained in Spanish (Bojanowski et al., 2017). Word embeddings are vector representations of words in a way that semantically related words have similar vectors. We employed FastText to measure the similarity between the question asked and the entire conversation. Specifically, our method calculates the similarity between all words in the question and all words in the conversation, with the exclusion of stop words and small words. This process generates a distribution of similarity values from which we take its mean

$$S = \frac{1}{n \cdot m} \sum_{i=1}^n \sum_{j=1}^m \left(\frac{\vec{V}_i \cdot \vec{W}_j}{\|\vec{V}_i\| \|\vec{W}_j\|} \right) \quad [1]$$

Non-parametric normalization.

The distributions of estimates for each question were centered on different values. To normalize these distributions, we used a non-parametric approach, originally developed in the outlier detection literature, and used in previous studies examining the wisdom of crowds (Barrera-Lemarchand et al., 2024; Navajas et al., 2018). This procedure involves subtracting the median value of the distribution and then dividing by the median absolute deviance, leading to a non-parametric z-score value. The rationale for normalizing our data was twofold. First, we used this procedure to reject outliers in the distribution of responses. Following previous studies, we discarded all responses that deviated from the median by more than 2.5 times the median absolute deviance. The second purpose of normalization was to average our results across different questions. This helps the visualization of our data, but our main findings can be replicated without any normalization (Table S4).

Robustness checks against answer look-up

All studies were conducted online. Although participants were instructed not to search for answers during the task, the online setting does not technically prevent such behavior. We therefore conducted additional robustness checks to evaluate whether the results could be driven by participants looking up the correct answers.

First, we examined how often group conversations produced the exact true value of the target quantity. If participants had systematically searched for the answers online, we would expect exact matches to occur frequently. In practice, only a small fraction of conversations produced answers that matched the true value exactly (between 0.6% and 2.8% across studies). For some questions, such matches could plausibly occur without looking up the answer because the correct value falls within commonly guessed ranges (e.g., “How many kilograms of ice cream does a person in Argentina eat per year?”). Therefore,

to address this concern conservatively, we repeated the main analyses after excluding all conversations that produced the exact correct answer. The results remain qualitatively unchanged (e.g., Mann–Whitney U test between “Fermi” and “Numbers” conditions in Study 2: $z = 1.95$, $p = 0.04$; Cohen’s $d = 0.25$).

Evaluating variability between human raters

We studied if the relationship between Fermi ratings and collective accuracy could reflect variability in how annotators interpreted the coding scheme. To evaluate this possibility, we examined the variance of Fermi ratings across annotators at the level of individual conversations. Disagreement among raters was generally small (mean $SD \approx 0.16$ Likert points), indicating high consistency in the evaluations. Moreover, between-rater variance did not differ significantly across experimental conditions (Mann–Whitney U test: $z = 1.8$, $p = 0.07$; Cohen’s $d = 0.29$) and did not correlate with either the linguistic markers of Fermi reasoning or collective estimation error (problem decomposition: $r = 0.07$, $p = 0.36$; guesstimation: $r = 0.09$, $p = 0.22$; collective error: $r = 0.02$, $p = 0.78$). As an additional robustness check, we repeated the main analyses after excluding conversations whose between-rater variance fell within the highest quartile of the distribution. The results remain unchanged under this restriction (Pearson correlation between Fermi ratings and collective error: $r = -0.08$, $p = 0.04$), indicating that the observed relationship between Fermi ratings and estimation accuracy is not driven by conversations for which annotators showed greater disagreement.

Linguistic markers

To quantify the extent to which conversations implemented specific reasoning processes, we developed a set of theory-driven linguistic indicators derived from the formal definition of the Fermi method and from the hypothesis of error-correction dynamics introduced in the main text. In all cases, the indicators were computed at the level of the

conversation by quantifying the prevalence of stylized expressions associated with each target process. Text was preprocessed in a case-insensitive manner, and accent marks did not affect matching.

To capture problem decomposition, we identified expressions typically used to break a problem into multiplicative or per-unit subcomponents. These included constructions involving the Spanish preposition “*por*” (which means “*by*”, “*per*”, or “*times*”, depending on context) followed by a noun or number, multiplicative operators (e.g., “*x*”, “***”) followed by a noun or number, and expressions based on the word “*cada*” (which translates to “*each*” or “*every*”) followed by a noun or number.

To capture guesstimation (i.e., approximate numerical reasoning under uncertainty), we identified expressions typically used to propose rough estimates. These included numerical ranges introduced by “*entre*” (meaning “*between*”) followed by a number, or modifiers indicating numerical uncertainty such as “*aproximadamente*” (or the colloquial “*aprox*”, meaning “*approximately*”) or “*alrededor*” (“*about*”) followed by a number.

To capture the explicit evaluation or revision of quantities, we identified expressions associated with magnitude adjustment and correction. These included forms of the verb *ser* (“to be”) combined with *mucho* (“a lot” or “too much”), *poco* (“a little” or “too little”), or *demasiado* (“too much”), as well as evaluative correction terms such as *mejor* (“better”).

LLM-based ratings

To obtain AI-based Fermi ratings, we queried each conversation using the application programming interface (API) of five large language models: gpt-4.1-mini, gpt-5-mini, gpt-5.2, claude-haiku-4.5, and claude-sonnet-4.6. Each model received the same prompt, which presented the formal definition of the Fermi method and asked the model to

evaluate the extent to which the conversation implemented that reasoning procedure on a 0–10 scale. Responses were constrained to return a single numerical score only.

The exact prompt was:

The Fermi Method consists of breaking the problem into simpler components, making rough estimations for each component, and calculating the final answer based on those estimations. You are an expert on the Fermi Method. Based on the definition above, evaluate the extent to which the participants in the following conversation used the Fermi Method. Rate the use of the Fermi Method on a scale from 0 to 10, where:

- 0 means the method was not used at all,*
- 10 means the method was fully and explicitly used throughout.*

Provide a single numerical score (0–10).

Return only the score (from 0–10), do not include any additional text.

Conversation: [conversation text]

All conversations were rated independently by each model using the same prompt and input format. The resulting scores were then used as AI-based Fermi ratings in the analyses reported in the main text.

Bayes factor analyses for equivalence testing

To assess evidence in favor of the absence of effects in several analyses, we complemented frequentist tests with Bayes factor analyses for equivalence testing. Bayes factors were computed using the default prior specification for standardized effect sizes proposed by Rouder et al. (2012). Specifically, we used a Cauchy prior distribution over standardized effect sizes centered at zero, following standard procedures in the Bayesian analysis of experimental designs. This default prior has desirable properties such as scale invariance and has become widely adopted in psychological research for evaluating evidence in favor of the null hypothesis.

Mixed-effects models.

To compare participants' estimates from the first and third stages in Study 1 (Fig. 1D), we used a generalized linear mixed-effects model, with random intercepts for each question, and fixed effects for each group and question.

The generalized linear model was:

$$\Delta E_{ij} = \beta_0 + \beta_1 \text{deliberated}_{ij} + \mu_j + \varepsilon_{ij} \quad [2]$$

where ΔE_{ij} denotes the change in error for observation i associated with question j and deliberated_{ij} is a binary predictor indicating whether the question was deliberated. The model includes a fixed intercept and fixed effect for deliberation, as well as a random intercept μ_j to account for repeated observations across questions. The same type of generalized linear mixed-effects models were used in Study 3 (Fig. 5C and Fig. 5D) with a dummy coding for the collective condition.

Code and data availability

All data and codes supporting our analyses are available at the Open Science Framework (https://osf.io/bru7k/?view_only=20b5fd8f315b4cb083186aeba294b632).

References

- Abay, S., & Filiz, S. B. (2020). Using Fermi problems to motivate 4th grade primary school students in math lessons. *Journal of Education, Theory and Practical Research*, 6(3), 268-277.
- Aceves, P., & Evans, J. A. (2024). Human languages with greater information density have higher communication speed but lower conversation breadth. *Nature Human Behaviour*, 8(4), 644-656.
- Albarracín, L., & Gorgorió, N. (2014). Devising a plan to solve Fermi problems involving large numbers. *Educational Studies in Mathematics*, 86(1), 79-96.
- Albarracín, L., & Gorgorió, N. (2019). Using large number estimation problems in primary education classrooms to introduce mathematical modelling. *International Journal of Innovation in Science and Mathematics Education*, 27(2).
- Albarracín, L., & Ärleback, J. B. (2025). Exploring the role of assumptions in mathematical modeling teacher training using Fermi problems. *ZDM—Mathematics Education*, 57(2), 213-228.
- Anderson, P. M., & Sherman, C. A. (n.d.). Applying the fermi estimation technique to business problems. *Journal of Applied Business & Economics*, 10(5).

- Bahrami, B., Olsen, K., Latham, P. E., Roepstorff, A., Rees, G., & Frith, C. D. (2010). Optimally Interacting Minds. *Science*, 329(5995), 1081–1085. <https://doi.org/10.1126/science.1185718>
- Barrera-Lemarchand, F., Balenzuela, P., Bahrami, B., Deroy, O., & Navajas, J. (2024). Promoting Erroneous Divergent Opinions Increases the Wisdom of Crowds. *Psychological Science*, 35(8), 872–886. <https://doi.org/10.1177/09567976241252138>
- Becker, J. A., Guilbeault, D., & Smith, E. B. (2022). The Crowd Classification Problem: Social Dynamics of Binary-Choice Accuracy. *Management Science*, 68(5), 3949–3965. <https://doi.org/10.1287/mnsc.2021.4127>
- Bermejo, V. J., Gago, A., Gálvez, R. H., & Harari, N. (2025). LLMs outperform outsourced human coders on complex textual analysis. *Scientific Reports*, 15(1), 40122.
- Bojanowski, P., Grave, E., Joulin, A., & Mikolov, T. (2017). Enriching word vectors with subword information. *Transactions of the Association for Computational Linguistics*, 5, 135–146.
- Carrillo, F., Cecchi, G. A., Sigman, M., & Fernandez Slezak, D. (2015). Fast distributed dynamics of semantic networks via social media. *Computational intelligence and neuroscience*, 2015(1), 712835.
- Dezecache, G., Dockendorff, M., Ferreiro, D. N., Deroy, O., & Bahrami, B. (2022). Democratic forecast: Small groups predict the future better than individuals and crowds. *Journal of Experimental Psychology: Applied*, 28(3), 525–537. <https://doi.org/10.1037/xap0000424>
- De Condorcet, M. M. J. A. N. (1785). *Essai sur l'application de l'analyse à la probabilité des décisions rendues à la pluralité des voix de Caritat*.
- Epp, D. A. (2017). Public policy and the wisdom of crowds. *Cognitive Systems Research*, 43, 53–61. <https://doi.org/10.1016/j.cogsys.2017.01.002>
- Espina Mairal, S., Bustos, F., Solovey, G., & Navajas, J. (2024). Interactive crowdsourcing to fact-check politicians. *Journal of Experimental Psychology: Applied*, 30(1), 3–15. <https://doi.org/10.1037/xap0000492>
- Frith, C. D. (2012). The role of metacognition in human social interactions. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 367(1599), 2213–2223.
- Galton, F. (1949). Vox Populi (1907) *Nature*, n. 1949, vol. 75, pp. 450–451 (traduzione di Romolo Giovanni Capuano\copyright). *Nature*, 75, 450–451.
- Gomilsek, T., Hoffrage, U., & Marewski, J. N. (2024). Fermian guesstimation can boost the wisdom-of-the-inner-crowd. *Scientific Reports*, 14(1), 5014.
- Gürçay, B., Mellers, B. A., & Baron, J. (2015). The Power of Social Influence on Estimation Accuracy. *Journal of Behavioral Decision Making*, 28(3), 250–261. <https://doi.org/10.1002/bdm.1843>
- Gutiérrez, L., & Keith, B. (2019). A Systematic Literature Review on Word Embeddings. In J. Mejia, M. Muñoz, Á. Rocha, A. Peña, & M. Pérez-Cisneros (Eds.), *Trends and Applications in Software Engineering* (Vol. 865, pp. 132–141). Springer International Publishing. https://doi.org/10.1007/978-3-030-01171-0_12
- Juni, M. Z., & Eckstein, M. P. (2015). Flexible human collective wisdom. *Journal of Experimental Psychology: Human Perception and Performance*, 41(6), 1588.
- Kameda, T., Toyokawa, W., & Tindale, R. S. (2022). Information aggregation and collective intelligence beyond the wisdom of crowds. *Nature Reviews Psychology*, 1(6), 345–357.

- Keller, A., Gerkin, R. C., Guan, Y., Dhurandhar, A., Turu, G., Szalai, B., Mainland, J. D., Ihara, Y., Yu, C. W., Wolfinger, R., Vens, C., Schietgat, L., De Grave, K., Norel, R., DREAM Olfaction Prediction Consortium, Stolovitzky, G., Cecchi, G. A., Vosshall, L. B., & Meyer, P. (2017). Predicting human olfactory perception from chemical features of odor molecules. *Science*, 355(6327), 820–826. <https://doi.org/10.1126/science.aal2014>
- Krockow, E. M., Kurvers, R. H. J. M., Herzog, S. M., Kämmer, J. E., Hamilton, R. A., Thilly, N., Macheda, G., & Pulcini, C. (2020). Harnessing the wisdom of crowds can improve guideline compliance of antibiotic prescribers and support antimicrobial stewardship. *Scientific Reports*, 10(1), 18782. <https://doi.org/10.1038/s41598-020-75063-z>
- Kurvers, R. H. J. M., Herzog, S. M., Hertwig, R., Krause, J., Carney, P. A., Bogart, A., Argenziano, G., Zalaudek, I., & Wolf, M. (2016). Boosting medical diagnostics by pooling independent judgments. *Proceedings of the National Academy of Sciences*, 113(31), 8777–8782. <https://doi.org/10.1073/pnas.1601827113>
- Landemore, H. E. (2017). Deliberative Democracy as Open, Not (Just) Representative Democracy. *Daedalus*, 146(3), 51–63. https://doi.org/10.1162/DAED_a_00446
- Landemore, H. E. (2012). Why the Many Are Smarter than the Few and Why It Matters. *Journal of Deliberative Democracy*, 8(1), Article 1. <https://doi.org/10.16997/jdd.129>
- Landemore, H. E., & Mercier, H. (2010). *‘Talking it Out’: Deliberation with Others Versus Deliberation Within* (SSRN Scholarly Paper No. 1660695). Social Science Research Network. <https://doi.org/10.2139/ssrn.1660695>
- Landemore, H., & Page, S. E. (2015). Deliberation and disagreement: Problem solving, prediction, and positive dissensus. *Politics, Philosophy & Economics*, 14(3), 229–254. <https://doi.org/10.1177/1470594X14544284>
- Lee, J., Li, T., & Shin, D. (2022). The Wisdom of Crowds in FinTech: Evidence from Initial Coin Offerings. *The Review of Corporate Finance Studies*, 11(1), 1–46. <https://doi.org/10.1093/rcfs/cfab014>
- Lorenz, J., Rauhut, H., Schweitzer, F., & Helbing, D. (2011). How social influence can undermine the wisdom of crowd effect. *Proceedings of the National Academy of Sciences*, 108(22), 9020–9025. <https://doi.org/10.1073/pnas.1008636108>
- Madirolas, G., & de Polavieja, G. G. (2015). Improving collective estimations using resistance to social influence. *PLoS Computational Biology*, 11(11), e1004594.
- Manning, C., & Schutze, H. (1999). *Foundations of statistical natural language processing*. MIT press.
- Mellers, B., Ungar, L., Baron, J., Ramos, J., Gurcay, B., Fincher, K., Scott, S. E., Moore, D., Atanasov, P., Swift, S. A., Murray, T., Stone, E., & Tetlock, P. E. (2014a). Psychological Strategies for Winning a Geopolitical Forecasting Tournament. *Psychological Science*, 25(5), 1106–1115. <https://doi.org/10.1177/0956797614524255>
- Mellers, B., Ungar, L., Baron, J., Ramos, J., Gurcay, B., Fincher, K., Scott, S. E., Moore, D., Atanasov, P., Swift, S. A., Murray, T., Stone, E., & Tetlock, P. E. (2014b). Psychological Strategies for Winning a Geopolitical Forecasting Tournament. *Psychological Science*, 25(5), 1106–1115. <https://doi.org/10.1177/0956797614524255>
- Moshman, D., & Geil, M. (1998). Collaborative Reasoning: Evidence for Collective Rationality. *Thinking & Reasoning*, 4(3), 231–248. <https://doi.org/10.1080/135467898394148>

Supplementary Information

Albarracín, L., & Gorgorió, N. (2014).	(The Fermi method is) "...based on designing a problem-solving strategy adapted to the problem which consists of breaking it into subproblems that may be solved separately by means of their respective estimations."
Albarracín, L., & Gorgorió, N. (2019).	"The procedure proposed by Fermi to his students was to decompose the original problem into simpler sub-problems to reach a solution of the original question by means of reasonable estimates or educated guesses"
Abay, S., & Filiz, S. B. (2020).	"Fermi problems are open-ended, non-routine problems that require students to make systematic guesses by making assumptions before starting on a solution using simple calculations."
Gomilsek et al. (2024)	"Fermian strategies prescribe decomposing an estimation problem into subtasks, solving the subtasks separately, and ultimately integrating those solutions into a final estimate."
Albarracín, L., & Ärleback, J. B. (2025)	"The Fermi estimation method is the classic approach to solving Fermi problems, in which assumptions are made to decompose the original problem into simpler subproblems that are solved separately using reasoned estimates based on knowledge of the real world."

Table S1: Definitions of Fermi problems or procedures in previous literature. Despite differences in terminology, these definitions converge on two core reasoning steps: problem decomposition and "guesstimation" of intermediate values which are then recombined to produce a final estimate.

Question	Correct Answer
How many goals were scored in total at the 2014 FIFA World Cup played in Brazil?	171
How many steps are there on the internal staircase of The Obelisk of Buenos Aires?	206
How many people (in millions) live in South America?	373
What is the total length (in meters) of line A of the Buenos Aires underground network?	10800
Until 2021, in how many Copa Libertadores finals has an Argentine team participated in?	37
How many kilograms of ice cream does a person in Argentina eat per year?	7
How many episodes are there in the American television series "Friends"?	236
How many countries have territory in the southern hemisphere, excluding the Antarctic continent?	48

Table S2: Questions tested in all experiments and their correct answers

Ind.	Accuracy			Confidence		
	TOST ($d_1 = 20\%$, $d_2 = 30\%$)		Bayes Factor	TOST ($d_1 = 20\%$, $d_2 = 30\%$)		Bayes Factor
1	$p_1 (x < d_1)$	$p < 10^{-200}$	15.0	$p_1 (x < d_1)$	4×10^{-5}	0.005
	$p_2 (x > d_2)$	10^{-8}		$p_2 (x > d_2)$	10^{-32}	
	C.I.	[0,0.022]		C.I.	[-0.036,-0.016]	
2	$p_1 (x < d_1)$	10^{-16}	52.7	$p_1 (x < d_1)$	10^{-11}	18.5
	$p_2 (x > d_2)$	2×10^{-14}		$p_2 (x > d_2)$	3×10^{-21}	
	C.I.	[-0.008,0.013]		C.I.	[-0.019,0.001]	
3	$p_1 (x < d_1)$	2×10^{-13}	48.7	$p_1 (x < d_1)$	$p < 10^{-200}$	3.90
	$p_2 (x > d_2)$	8×10^{-17}		$p_2 (x > d_2)$	4×10^{-9}	
	C.I.	[-0.014,0.007]		C.I.	[0.004,0.024]	
4	$p_1 (x < d_1)$	6×10^{-12}	12.2	$p_1 (x < d_1)$	$p < 10^{-200}$	0.34
	$p_2 (x > d_2)$	2×10^{-23}		$p_2 (x > d_2)$	5×10^{-6}	
	C.I.	[-0.020,0.001]		C.I.	[0.010,0.032]	

Table S2: Results from two one-sided tests (TOST) of equivalence and Bayes factor analyses for the proportion of interventions of each individual (1-4), ordered by accuracy (first columns) and confidence (last columns). All TOSTs indicate that all proportions of interventions are within the 20-30% interval.

	Humans	GPT 4.1 mini	GPT 5.1 mini	GPT 5.2	Claude Haiku 4.5
GPT 4.1 mini	0.80				
GPT 5.1 mini	0.70				
GPT 5.2	0.78				
Claude Haiku 4.5	0.84				
Claude Sonnet 4.6	0.83	0.88	0.80	0.86	0.86

Table S3: Pearson correlation coefficient between Fermi ratings produced by five different large language models (LLMs). All comparisons are statistically significant with $p < 0.001$.

	β	SE	t	p-val	95% C.I.
Intercept	0.20	0.02	8.00	2×10^{-14}	[0.15 ; 0.24]
Fermi Ratings	-0.01	0.003	-2.98	0.003	[-0.02 ; -0.003]
DF	323				
Log-likelihood	107				
AIC	-206				
BIC	-191				

Table S4: Replication of the main result of Study 1 without normalization. Generalized linear mixed-effect model of Collective Error, with the Fermi Ratings as a fixed effect, and random intercepts for each question.

Study 1 (N=500)

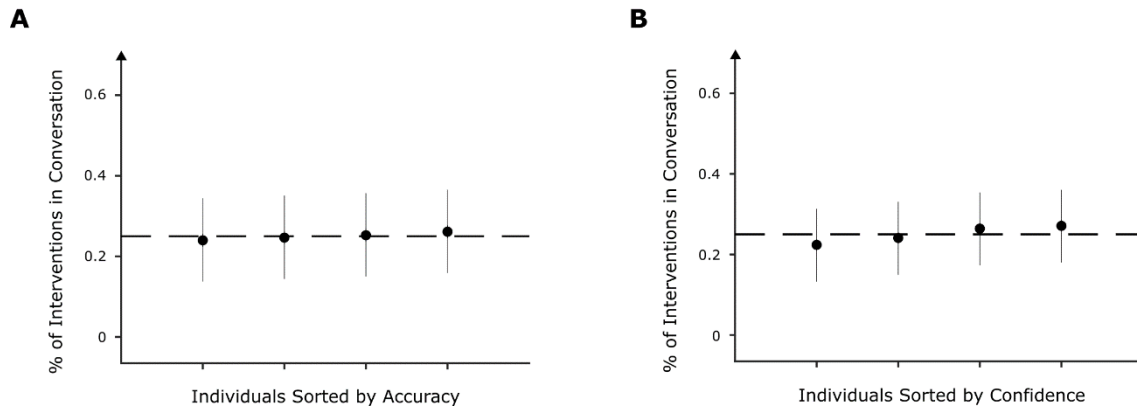


Fig. S1. Accuracy and confidence for each individual for Experiment 1. (A) Percentage of interventions (messages) of each individual for each conversation. Individuals are sorted by accuracy. The full lines represent the standard deviation of the mean, and the dotted line represents 25% **(B)** Percentage of interventions (messages) of each individual for each conversation. Individuals are sorted by confidence. The full lines represent the standard deviation of the mean, and the dotted line represents 25%.

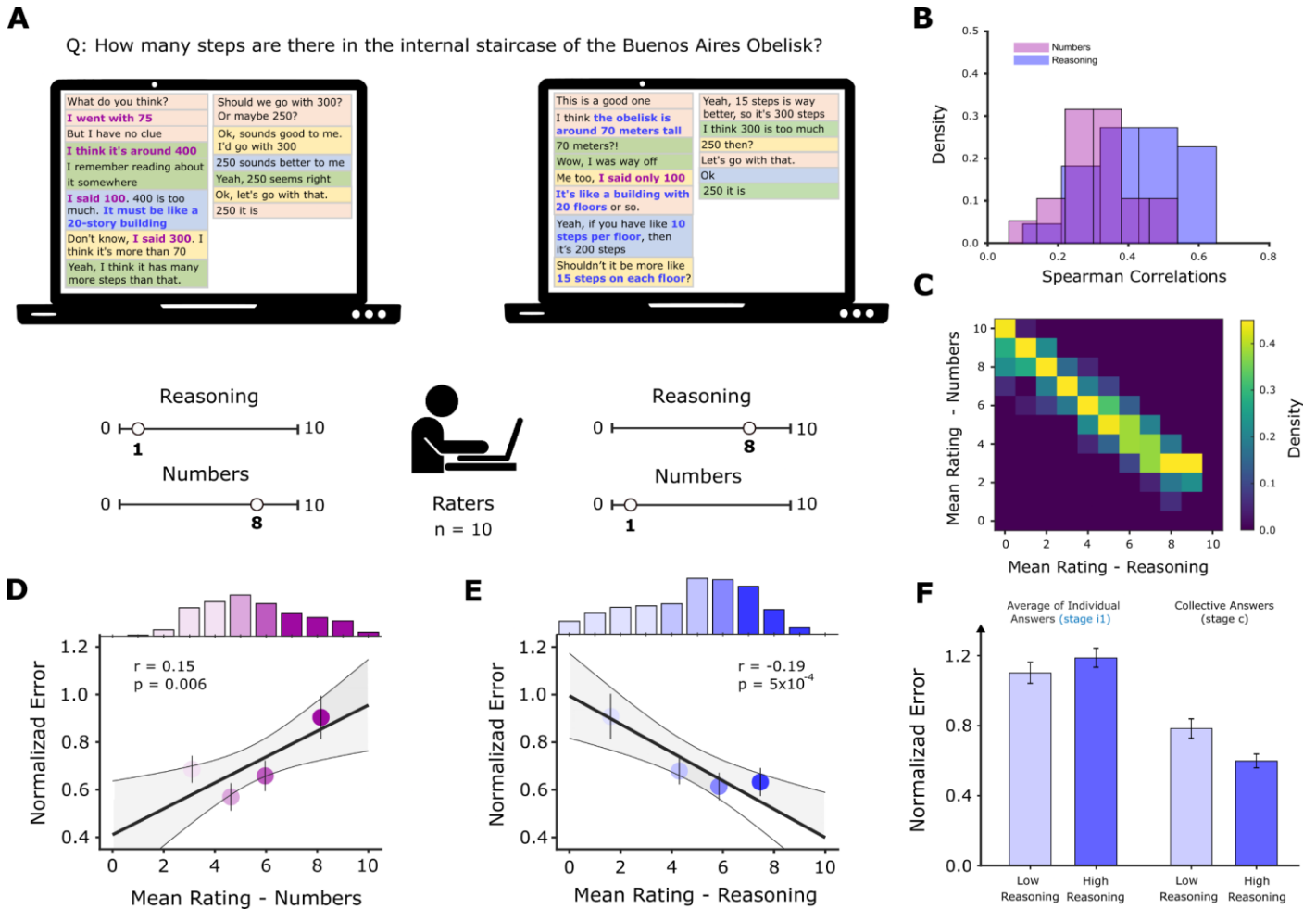


Fig. S2. Analysis of Conversations. Replication of Figure 2 under a different way of coding conversations, (A) Example of conversations that lead to opposing ratings of the strategies used in them. 6 raters were trained on how to classify conversations in terms of whether participants pooled their previous individual answers (“Numbers strategy”) or reached a new collective answer by employing a reasoning procedure (“Reasoning strategy”). **(B)** Distribution of correlations between the values provided by each rater with each other, for each strategy. The correlation of the values provided by different raters is generally high, for both strategies. **(C)** Density of ratings for the Reasoning and Numbers Strategies. Ratings for both strategies were strongly negatively correlated: whenever raters on average assigned a high value for the Reasoning strategy in a conversation, they also tended to assign on average a very low value for the Numbers strategy to that conversation. **(D)** Normalized error as a function of the “Numbers strategy” ratings (averaged across raters). We plot the quartiles of the distribution of normalized errors, along with the best linear fit of the data, and the 95% confidence intervals. The bars above the main graph show the actual distribution of normalized errors. We also show the Pearson correlation coefficient, and the corresponding p-value **(E)** Normalized error as a function of the

“Reasoning strategy” ratings (averaged across raters). We plot the quartiles of the distribution of normalized errors, along with the best linear fit of the data, and the 95% confidence intervals. The bars above the main graph show the actual distribution of normalized errors. We also show the Pearson correlation coefficient, and the corresponding p-value. **(F)** Variation of Normalized Error between groups and averages of the individual answers. We categorize a group for each conversation as either “low Reasoning” (when the Reasoning rating was lower than the Numbers rating) or “high Reasoning” (when the Reasoning rating was higher than the Numbers rating).

Study 2 (N=240)

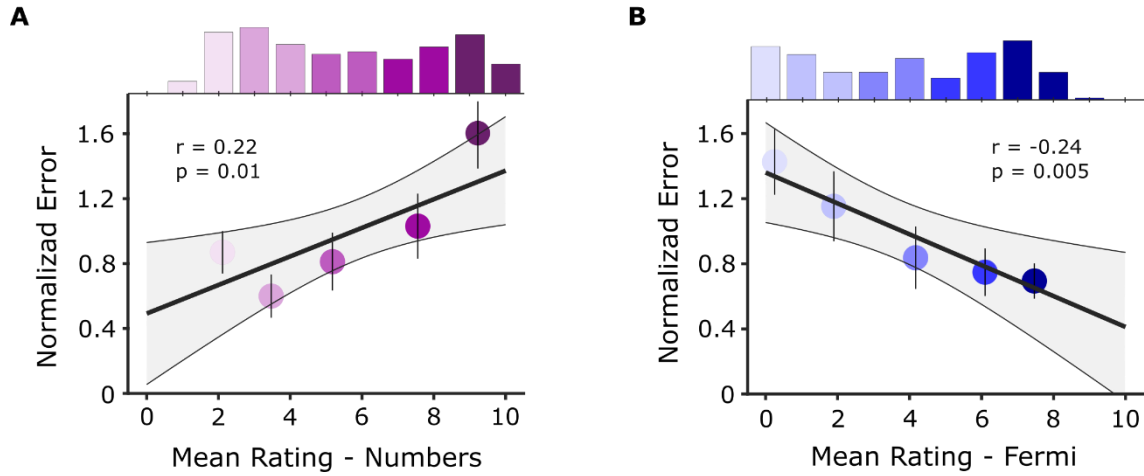


Fig. S3. Normalized error as a function of human ratings. Replication of Figs 2D and 2E (Study 1) in Study 2. (A) Normalized error as a function of the “Numbers strategy” ratings (averaged across raters). We plot the quintiles of the distribution of normalized errors, along with the best linear fit of the data, and the 95% confidence intervals. The bars above the main graph show the actual distribution of normalized errors. We also show the Pearson correlation coefficient, and the corresponding p-value **(B)** Normalized error as a function of the “Fermi strategy” ratings (averaged across raters). We plot the quintiles of the distribution of normalized errors, along with the best linear fit of the data, and the 95% confidence intervals. The bars above the main graph show the actual distribution of normalized errors. We also show the Pearson correlation coefficient, and the corresponding p-value

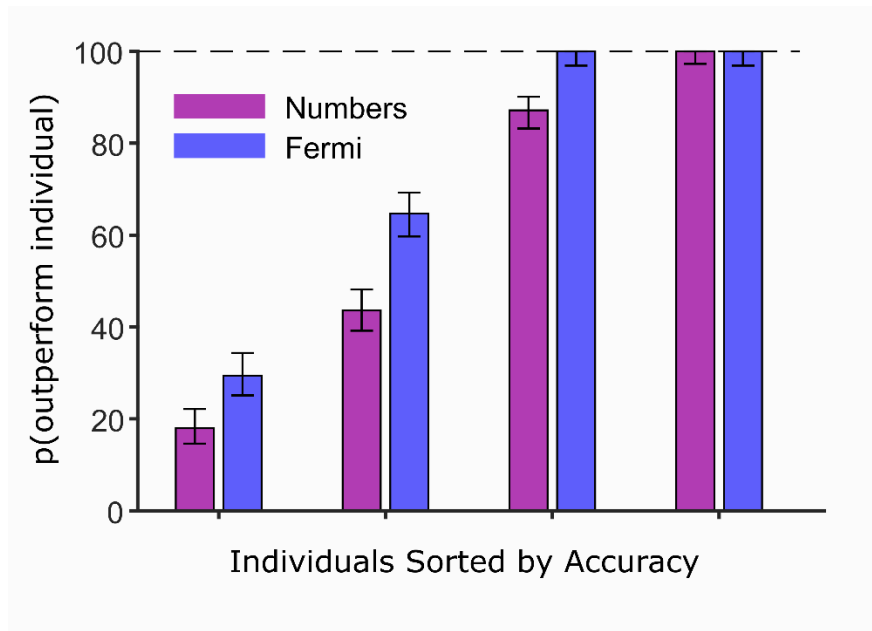


Fig. S4. Probability that the collective estimate outperforms individual estimates. In Study 2, individuals within each group were ranked according to their individual estimation error (from most to least accurate). The figure shows the probability that the collective estimate outperforms the estimate of each individual in the group under the two experimental conditions (“Numbers” and “Fermi”). Error bars represent ± 1 s.e.p.

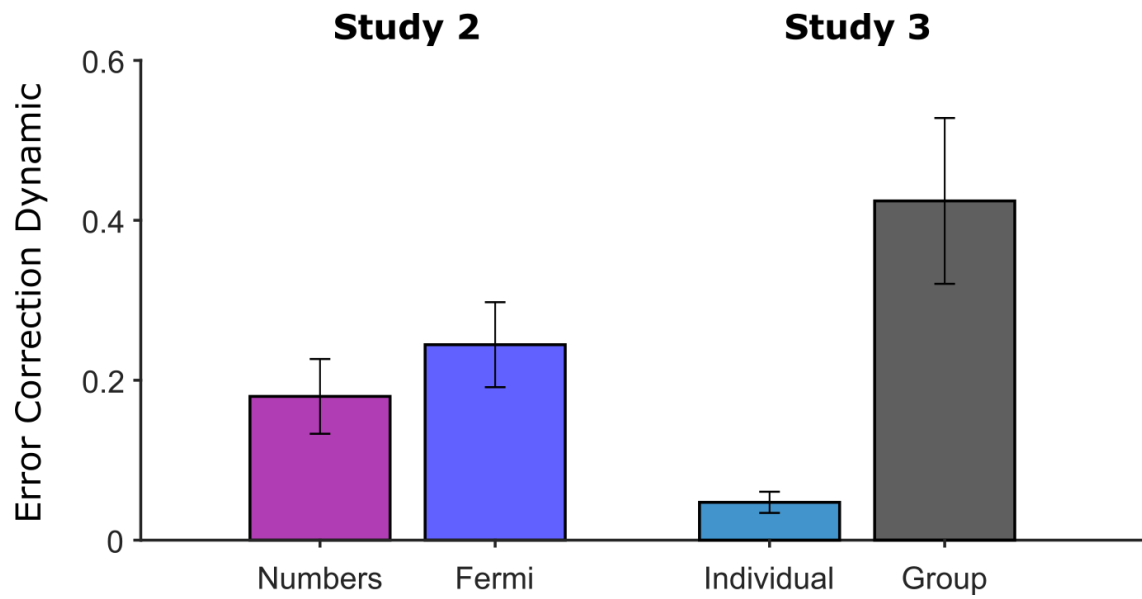


Fig. S5. Error correction dynamics in Studies 2 and 3. In Study 2, the presence of linguistic markers reflecting error correction dynamics (e.g., “too little”, “too much”, “better”) was not statistically different between conditions involving different instructions. In Study 3, the same linguistic marker led to a pronounced and significant difference between the individual and collective condition. Bars show the average value of the metric for each study and condition ± 1 s.e.m.

Supplementary Note 1

The instructions provided to the raters were the following:

Your task is to read conversations between groups of four individuals discussing factual topics (e.g., “What is the height of the Eiffel Tower?”). Each group was instructed to “try to reach consensus” on these topics. There are eight different questions in total.

All participants first completed an individual stage, where they provided answers to the eight questions on their own. Afterward, they were divided into groups and asked to discuss 4 of the 8 questions, with five minutes allotted per question (they could see their own previous answers). Most groups did reach consensus, but this was not required.

As you read each conversation, you will be asked to rate, on a scale from 0 (completely absent) to 10 (completely present), the extent to which two possible strategies for reaching consensus were used.

The first strategy involves combining the estimates from the initial individual stage. This could mean calculating an average, weighting responses by confidence, discarding some answers in favor of others, or any other approach that makes use of their original estimates. We call this the “Numbers Strategy.”

The second strategy involves breaking down the problem into smaller parts, estimating those quantities, and then combining them to arrive at a final answer. For example, for the question “What is the height of the Eiffel Tower?”, one might assume the Tower is similar to a 100-story building (first estimate), and if each story is 3 meters tall (second estimate), then the Tower would be about 300 meters high (which in this case happens to be correct, though this is not always so). We call this the “Fermi Strategy.”

Thus, for each conversation, you will be asked to answer two questions:

- a) Did they reach consensus by combining their initial estimates from the first stage? (Numbers Strategy) / 0 = completely absent, 10 = completely present
- b) Did they reach consensus by breaking the problem into smaller parts? (Fermi Strategy) / 0 = completely absent, 10 = completely present

Note that the ratings for a) and b) do not need to sum to 10. A group may have used both strategies, or neither. Also, be sure to use the full scale rather than only the extremes (0 and 10).