



Optimizing competence in the service of collaboration

Yang Xiang^{a,*}, Natalia Vélez^b, Samuel J. Gershman^{a,c,d}

^a Department of Psychology, Harvard University, United States of America

^b Department of Psychology, Princeton University, United States of America

^c Center for Brain Science, Harvard University, United States of America

^d Center for Brains, Minds, and Machines, MIT, United States of America

ARTICLE INFO

Dataset link: https://github.com/yyyyxiang/competence_training

Keywords:

Collaboration
Planning
Competence
Plasticity
Training
Teaching
Social cognition

ABSTRACT

In order to efficiently divide labor with others, it is important to understand what our collaborators can do (i.e., their *competence*). However, competence is not static—people get better at particular jobs the more often they perform them. This plasticity of competence creates a challenge for collaboration: For example, is it better to assign tasks to whoever is most competent now, or to the person who can be trained most efficiently “on-the-job”? We conducted four experiments ($N = 396$) that examine how people make decisions about whom to train (Experiments 1 and 3) and whom to recruit (Experiments 2 and 4) to a collaborative task, based on the simulated collaborators’ starting expertise, the training opportunities available, and the goal of the task. We found that participants’ decisions were best captured by a *planning* model that attempts to maximize the returns from collaboration while minimizing the costs of hiring and training individual collaborators. This planning model outperformed alternative models that based these decisions on the agents’ current competence, or on how much agents stood to improve in a single training step, without considering whether this training would enable agents to succeed at the task in the long run. Our findings suggest that people do not recruit and train collaborators based solely on their current competence, nor solely on the opportunities for their collaborators to improve. Instead, people use an intuitive theory of competence to balance the costs of hiring and training others against the benefits to the collaboration.

1. Introduction

Division of labor is central to collaboration: Even young children can appropriately assign tasks based on the relative competence of their collaborators (Baer & Odic, 2022; Magid et al., 2018), or take their collaborator’s physical constraints into account when planning their own actions (Warneken et al., 2014). One of the chief advantages of division of labor is that it allows individual workers to get better at specialized tasks through practice. In *The Wealth of Nations*, Smith (1937) gives the example of a contest between a teenager and a smith: Through sheer practice, a teenager who solely makes nails, day in and day out, can vastly outperform a smith who only occasionally makes nails as one of the many products of her forge. Smith (1937) argued that these increases in efficiency from division of labor could directly account for extraordinary increases in the wealth of nations. Indeed, many of the benefits of collaboration and division of labor stem precisely from the fact that competence is not static.

Taking into account how people’s competence changes with practice is particularly important when making decisions about whom to recruit for a collaborative task, or how much to invest in their training. For example, a company might hire a less

* Correspondence to: 52 Oxford St., Cambridge MA 02138, United States of America.

E-mail address: yyx@g.harvard.edu (Y. Xiang).

experienced candidate who – with appropriate training – could qualify for the job, as opposed to a highly-skilled candidate who requires a hefty salary. Similarly, graduate school advisors might admit a student who shows great potential, given opportunities to train and learn. Training also allows current team members to become better suited for their roles. Organizations investing in effective training programs that enhance employees' capabilities tend to achieve both short- and long-term benefits for both the employees and the organizations (Nda & Fard, 2013). Effective training fills the gap between desired performance and actual employee performance; by improving employees' performance and capabilities, companies increase their organizational productivity (Elnaga & Imran, 2013). Formal employee training programs are also unique in their ability to bring below-average firms up to the performance level of comparable businesses (Bartel, 1994). In short, people often make decisions in real-world collaborations by considering not only how capable people are already, but also how capable they *will become* over the course of a long-term collaboration. However, little is known about the lay intuitions behind how people make decisions about whom to recruit, whom to train, and how to train them.

Answers to these questions hinge on many factors, including what we know about the team structure, what training resources are available, how much time we have for training, etc., and also require that we take into account how our collaborators' competence may change on the job. These decisions can be characterized as a sequential decision problem that requires long-term prospection. Recent work suggests that – much as people can use *planning* to solve sequential decision problems (Daw & Dayan, 2014; Huys et al., 2015) in non-social settings – they can also plan joint actions efficiently (Curioni, 2022; Török et al., 2019, 2021) and use intuitive theories of other people's minds to plan interventions on their mental states (Ho et al., 2022), including beliefs and desires (Baker et al., 2017, 2009; Jara-Ettinger et al., 2016; Premack & Woodruff, 1978), and also representations of competence and effort (Xiang et al., 2023b). Thus, one possibility is that people may make decisions about whom to recruit and train for a collaborative task by anticipating how others' competence may change over time and planning accordingly. However, planning in general is computationally expensive, and these costs may be prohibitive in the context of collaboration, where optimal planning requires recursive theory of mind (Vélez & Gweon, 2021).

Thus, people may instead use simpler heuristics to decide whom to assign to what task without planning ahead (Dhami, 2003; Payne et al., 1993, 1996). For example, people may simply assign tasks based on who is more competent now, without considering how their competence will change as a result of training. Indeed, past work has found that adults tend to select teams whose combined expertise suffices for a particular job (Xiang et al., 2023b), and that children tend to use relative differences in ability to assign people to tasks (Baer & Odic, 2022; Magid et al., 2018). Another possibility is that – rather than treating training as a means to an end – people might instead train agents who need it the most. If this is the case, then they may selectively train the weakest agent, or the agent who stands to improve the most from a single bout of training, without planning ahead to consider what the group's overall competence will be after training.

Across four experiments ($N = 396$), we show that people engage in planning when they make training and hiring decisions regarding simulated agents. The first two experiments are a collaborative box-lifting task framed as a game show. In Experiment 1, participants trained agents to prepare them to lift a heavy box together. In Experiment 2, participants first selected two agents out of a pool of four candidates, then trained the selected agents to lift a heavy box together. Experiments 3 and 4 generalized the task to a more education-oriented context, where participants recruited and trained students for a math Olympiad. We compared the human judgments to a *planning* model and four alternative models that base training decisions on heuristics: an *exploitation* model that selects the strongest agent, an *equity* model that selects the weakest agent, a *learning* model that trains agents who will improve the most after a single training step, and an *equality* model that invests equally in every agent's training. We found that the planning model qualitatively and quantitatively provided the best match to the data.¹ In the following section, we describe these models in detail.

2. Theoretical framework

In Experiments 1 and 2, pairs of simulated agents play a contest where they attempt to lift a heavy box (i.e., the target) together to win a prize. Before the contest, participants can increase the agents' strength by assigning them boxes to train with on their own. Experiments 3 and 4 apply the same framework to a similar problem with a different framing: Participants – imagining themselves as high school math coaches – recruit and train student contestants for a math Olympiad. Participants can increase students' math abilities by assigning them problem sets to complete. Since these two paradigms differ only in the framing, we describe below how we set up the models for the collaborative box-lifting task in Experiments 1 and 2. We start by describing how individual “trainees” (i.e., simulated agents) gain strength, based on the effort they invest into the lift, and then describe a suite of models of how “trainers” (i.e., participants) decide whom to train and how to train them.

2.1. Modeling trainees' behavior

Suppose each agent a has a strength S_a and exerts effort $E_a \in [0, 1]$ to lift a training box of weight W . Effort E_a defines the proportion of an agent's strength that is applied to lifting.² Therefore, agent a applies a force of $S_a \cdot E_a$ to lift, and succeeds if this

¹ Our data and code are publicly available at https://github.com/yyxiang/competence_training.

² For computational tractability, we constrained the effort space to discrete values, ranging from 0 to 1 with increments of 0.01.

force equals or exceeds the box weight W (i.e., $S_a \cdot E_a \geq W$). The agent's utility function is defined as the reward from lifting the box minus the effort cost (Xiang et al., 2023b):

$$U(E) = R \cdot L - C(E), \quad (1)$$

where R is the reward the agent receives from successfully lifting the box, $C(E)$ is the effort cost, and $L = \mathbb{I}[E_a \cdot S_a \geq W]$. The indicator function $\mathbb{I}[\cdot] = 1$ if its argument is true, 0 otherwise; thus, $L = 1$ indicates that the lift was successful. We model agents as deterministic utility maximizers (i.e., agents always choose the effort level that maximizes the utility function):

$$E_a^* = \operatorname{argmax}_{E_a} U(E_a). \quad (2)$$

Note that in this phase, where agents are lifting boxes individually, we do not yet need to consider collaborative optimization of effort.

For simplicity, we assume that agents are highly motivated to lift the box (i.e., $R \gg C(E)$). Therefore, E_a^* is effectively:

$$E_a^* = \begin{cases} W/S_a & \text{if } S_a \geq W \ (L = 1) \\ 0 & \text{if } S_a < W \ (L = 0) \end{cases} \quad (3)$$

This means that an observer can predict an agent's effort once they know the agent's strength and the weight of the box.

We make a simple yet plausible assumption about how strength increases with training: Strength increases in proportion to effort, which can be formalized by the following update rule:

$$\Delta S_a = \rho \cdot E_a, \quad (4)$$

where $\rho \in [0, 1]$ is a learning rate. Here, we set ρ to 1, such that strength increases by the level of effort exerted. If we assume that agents exert the optimal effort for the task (Eq. (3)), we see that failed lifts do not increase agents' strength – as the optimal effort is 0 – whereas a successful lift increases it in proportion to the effort exerted. Because agents' strength increases after each training bout, assigning the same weight a second time will require less effort and increase the agents' strength less, consistent with the “diminishing returns” principle of strength training (Morton et al., 2019).

2.2. Modeling trainers' decisions

In Experiment 1, participants trained pairs of agents so that they could lift a target box of weight W together. We define *training strategies* as the sequences of actions taken by participants during training; at each step, participants can choose to pay a small training cost to train one agent using a box selected from an array, or they can choose to not train any agents at all to save on training costs. We use S'_{ai} to denote agent a 's strength after training according to training strategy i . If the agents' combined strength equals or exceeds W by the end of training ($\sum_a S'_{ai} \geq W$), then the joint lift succeeds ($L = 1$), and participants receive a reward of R minus the cumulative training costs $K(i)$ of training strategy i . We set $K(i)$ as a small, fixed fee k for every unit of weight used during training; heavier boxes thus deduct higher costs from participants' bonus.

In Experiment 2, participants additionally chose which agents to recruit prior to training. Participants first selected two agents out of a pool of four candidates to form a team t , then trained the two selected agents ($a \in t$) to prepare the two of them to lift a heavy box together. The training process was identical to Experiment 1, although there was no cost associated with training (i.e., $K(i) = 0$). Instead, participants paid a hiring fee for each agent recruited on the team, which we set as a small, fixed fee h for every unit of strength. We denote the total fee for team t by $H(t)$.

To capture how participants make these decisions, we first describe a planning model that anticipates how agents' strengths will change as a result of the training strategy, then describe four alternative models that use heuristics to make decisions.

2.2.1. Planning model

For the *planning model*, the trainer's goal is to obtain the reward with minimal cost. Thus, the action value of training strategy i is defined as:

$$Q_i^{\text{planning}} = R \cdot L - K(i), \quad (5)$$

where L indicates whether the joint lift is successful by comparing the sum of agents' forces to the target box weight. L is defined as:

$$L = \mathbb{I} \left[\sum_a E'_a \cdot s'_{ai} \geq W \right], \quad (6)$$

where E'_a indicates the level of effort agent a exerts during the joint lifting.

Since we assume that agents are highly motivated to exert effort when necessary (i.e., $E'_a \approx 1$), the lift function becomes:

$$L \approx \mathbb{I} \left[\sum_a s'_{ai} \geq W \right] \quad (7)$$

In many cases, some strategies are equally preferred, and some others are only slightly less preferred. In light of this, we decided to add stochasticity to the choice process, instead of keeping the training strategy with the highest action value. We passed action

values through a softmax function so that the model chooses training strategies stochastically. Strategies with higher action values are more likely to be chosen. The probability of selecting training strategy i is:

$$P_i = \frac{e^{\gamma Q_i}}{\sum_{j=1}^K e^{\gamma Q_j}}, \quad (8)$$

where γ is the inverse temperature parameter that controls the stochasticity of the predictions by scaling the action values before applying softmax, and K is the total number of training strategies.

The prediction of agents' strength after training is therefore the expected value of the agent's updated strength—in other words, a weighted average of updated strengths from all the possible training strategies:

$$S'_a = \sum_i P_i \cdot S'_{ai} \quad (9)$$

In the team selection problem in Experiment 2, the action value for choosing each team is defined as:

$$Q_t^{\text{planning}} = R \cdot L - H(t), \quad (10)$$

where L is again an indicator function returning whether the agents lifted the target box together:

$$L = \mathbb{I} \left[\sum_{a \in t} S'_a \geq W \right] \quad (11)$$

Here S'_a indicates the selected agents' strengths post-training following the same calculation described above (Eq. (9)), except that we add the constraint that only agents selected on the team ($a \in t$) can receive training.

2.2.2. Alternative models

We compare the planning model to four alternative models that use heuristics in lieu of planning: an *exploitation model*, an *equity model*, a *learning model*, and an *equality model*. Same as above, we use S_a to denote agent a 's strength before training, and S'_a to denote agent a 's strength after training. Below, we describe how each of these models computes the action values of each strategy. The action values are then passed through the softmax function.

The **exploitation model** maximizes the strongest agent's strength after training while considering the costs of training and hiring. The action value of each training strategy is therefore:

$$Q_i^{\text{exploitation}} = \max_a S'_{ai} - K(i) \quad (12)$$

And the action value of selecting each team is:

$$Q_t^{\text{exploitation}} = \max_a S'_a - H(t) \quad (13)$$

Note that here and elsewhere, S'_a is implicitly dependent on the team t since only agents in t are trainable.

The **equity model** minimizes the strength difference between the strongest agent and the weakest agent after training while considering the costs of training and hiring:

$$Q_i^{\text{equity}} = \min_a S'_{ai} - \max_a S'_{ai} - K(i) \quad (14)$$

$$Q_t^{\text{equity}} = \min_a S'_{ai} - \max_a S'_a - H(t) \quad (15)$$

The **learning model** maximizes the total strength gain from training while considering the costs of training and hiring:

$$Q_i^{\text{learning}} = \sum_a S'_{ai} - S_a - K(i) \quad (16)$$

$$Q_t^{\text{learning}} = \sum_a S'_a - S_a - H(t) \quad (17)$$

The **equality model** invests equally in every agent's training by indiscriminately choosing training strategies and teams (inspired by Hertwig et al., 2002; Messick, 1993) while considering the costs of training and hiring:

$$Q_i^{\text{equality}} = -K(i) \quad (18)$$

$$Q_t^{\text{equality}} = -H(t) \quad (19)$$

We also considered a Rawlsian model that maximizes the strength of the weakest agent after training (Rawls, 2001) and a utilitarian-style model that maximizes the sum of the agents' strengths after training (Bentham, 1970; Mill, 2016). The Rawlsian model makes similar predictions as the equity model because maximizing the weakest agent's strength after training is essentially reducing the strength difference between the strongest and the weakest agents (for example, both models tend to focus entirely on training the weakest agent). The utilitarian-style model makes similar predictions as the learning model because maximizing the sum of the agents' strengths after training is equivalent to maximizing total strength gain. Since these two models make similar predictions as the ones we described above, we focus on the four aforementioned alternative models (exploitation, equity, learning, and equality) in the interest of parsimony.

In Experiments 3 and 4, we extend this framework into a math Olympiad task, where agents' strength is replaced with math abilities, and box weights are replaced with difficulties of math problems. Everything else carries over into the new task.

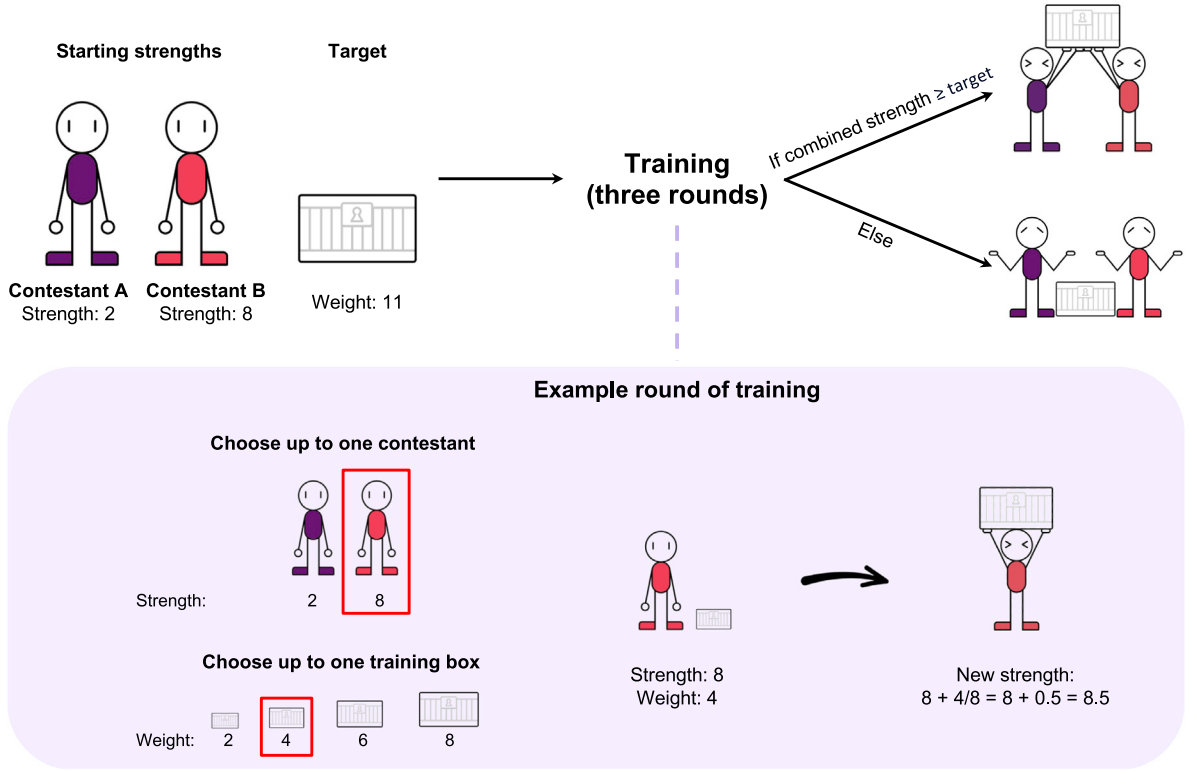


Fig. 1. Example contest in Experiment 1. Participants saw two agents of varying strengths and a target box. They had three rounds to train the agents. In each round of training, they could either assign a box to an agent or not train anyone. Agents gained strength equivalent to the level of effort they put forth. After training, agents attempted to lift the target box together.

2.3. Parameter fitting

Our theoretical framework contains a single free parameter: the inverse temperature parameter γ in Eq. (8). For every model, we fit this parameter to each participant's choices (training strategies in Experiments 1 and 3 and teams in Experiments 2 and 4) using maximum likelihood estimation through a grid search with search space ranging from 5 to 15, with increments of 0.5. We chose this range because, below the lower bound, the alternative models approach uniform distributions (specifically, exploitation, equity, and learning models in Experiments 1 and 3, and all four alternative models in Experiments 2 and 4).

3. Experiment 1

The purpose of Experiment 1 was to examine how participants train teammates in preparation for a collaborative task. Specifically, we asked whether participants' decisions are better captured by the planning model or by alternative heuristic models. We designed four scenarios in which the models described above make qualitatively distinct predictions. The planning model predicts that people will train agents to be *just strong enough* to clear the task, while the exploitation, equity, and learning models will try to maximize agents' final strength by maximizing the strongest agent's strength, reduction in the strength difference between the strongest and weakest agents, and the total strength gain, respectively, after accounting for training costs. The equality model will essentially only consider training costs since, otherwise, it prefers every training strategy equally. To ground the scenarios in a realistic setting, we framed the experiment as a game show consisting of four contests.

3.1. Materials and methods

3.1.1. Participants

We recruited 99 participants via Amazon's Mechanical Turk platform (MTurk).³ In this and subsequent experiments, to ensure data quality, our task was only available to MTurk workers who have had at least 1,000 studies approved and an approval rate higher

³ In Experiments 1, 3, and 4, we stopped data collection once 100 participants submitted their responses on MTurk, as stated in the pre-registration, but only received 99 participants' data due to a server error.

than 98%, and participants were only allowed to proceed to the experiment after answering all questions in a comprehension check correctly.

Participants received a base pay of \$1 and a potential bonus payment of up to \$4. The bonus payment depended on the training boxes participants used to train the contestants and whether the lifting turned out successful. Specifically, for every successful lift, participants received \$1; for every weight unit they used, \$0.04 was deducted from their bonus payment, regardless of the lift outcome. The experiment was approved by the Harvard Institutional Review Board and pre-registered at https://aspredicted.org/1GF_TVQ.

3.1.2. Procedure

Fig. 1 shows the task setup. The game show consisted of four contests, each with a different pair of contestants. In each contest, participants first saw the contestants' starting strengths and the weight of a heavy box, and were asked to train contestants to prepare them to lift the heavy box together. In order to make the strengths and weights intuitive to the participants, we expressed them using the same units. Each contestant's strength determines the heaviest weight they can lift given an all-out effort; a contestant with a strength of 5 would fail to lift boxes with a weight of 6 or higher on their own, no matter how much effort they exerted.

Participants were provided with four training boxes in all contests (Weights 2, 4, 6, and 8) and were told that contestants would gain strength based on the effort they put into lifting each box. For example, if a contestant with a starting strength of 6 was assigned a training box of weight 3, they would only need to put in a 50% effort to lift the box; consequently, the contestant's strength would increase by 0.5. However, if the box is too heavy for them to lift, contestants will give up and not gain any strength. To ensure that participants understood the task setup, they completed three practice trials where they observed how a contestant's strength changed by lifting different boxes.

In each contest, participants had up to three turns to train the contestants. On each turn, they could choose which agent to train and which training box to assign to them, or they could choose to not train anyone at all. Participants saw the updated strengths of the contestants after each turn. After training, participants saw the two contestants attempt to lift the heavy box together; the contestants were able to do so only if their combined strength was equal to or greater than the weight of the heavy box. Links to the experiments can be found at https://github.com/yyxiang/competence_training.

3.2. Results

Fig. 2A shows average behavioral data and model predictions in a series of 2-D plots. Each subplot corresponds to a different scenario, and each line shows how both agents' strength changed, on average, over the course of training in that scenario; the weaker agent's strength is plotted on the x -axis, and the stronger agent's strength is plotted on the y -axis. Each point in the participants' responses (Data) denotes the average strength of the two agents in each round of training. Each point in the model predictions denotes the weighted average strength calculated from Eq. (9), using values of γ fit to individual participants. Because agents' strengths can only increase or stay the same, their starting strengths are at the origin of each plot, and points along a line that are farther from the origin occurred later in training. The dashed black lines indicate the minimum total strength required to lift the target box. In other words, agents are strong enough to complete the task once their combined strength passes the dashed line. From these plots, we can see that participants' average responses aligned closely with the predictions of the planning model, and they both have endpoints close to the dashed line; both participants and the planning model tended to train agents to be just strong enough for the task, weighing rewards and costs. By contrast, the alternative models failed to capture this pattern: The exploitation model trained the stronger agent as much as possible, the equity model trained the weaker agent as much as possible, and the learning model trained the two agents for the largest strength gain possible, all without considering the target weight. The equality model, on the other hand, did not train agents to be strong enough to lift the box in three out of four scenarios.

To quantitatively assess how well each model captures the data, we computed correlations between the participant-averaged data and model predictions of agents' strength at every step of training. The planning model and the equality model have the highest correlation with the data (Pearson's $r = 0.999, p < .0001$ for both), followed by the learning model ($r = 0.997, p < .0001$), exploitation model ($r = 0.992, p < .0001$), and equity model ($r = 0.983, p < .0001$). Note that these correlation coefficients are all very high because agents' strength can only increase or stay the same with training, and because agents' relative strength does not change much during scenarios, as the maximum strength increase possible during training is small compared to the difference in starting strengths between agents. Quantitatively comparing the average model-predicted strengths obscures differences between our models. Therefore, we deviated from our pre-registered plan and instead compared individual participants' training strategies to model predictions using random-effects Bayesian model selection (Rigoux et al., 2014). For each participant, we first evaluated the fit of each model using $BIC = k \ln(n) - 2 \ln(L)$, where k is the number of free parameters, n is the number of scenarios, and L is the maximized value of the log likelihood. We then approximated the log model evidence for each participant as $-0.5BIC$ (Bishop, 2006) and used it to compute the protected exceedance probability for each model, which can be interpreted as the probability that each model is the most frequently occurring in the population.

Overall, the planning model has a protected exceedance probability of approximately 1, which means that the planning model is the most likely model in the population. To compare how well each model fits individual participants' responses, Fig. 3A shows the distribution of model evidences for each model for each participant. The planning model had the highest model evidence. It is worth noting that Experiment 1 was not originally designed to discriminate between the planning model and the equality model; we differentiate between them in Experiment 3, below, with four new scenarios.

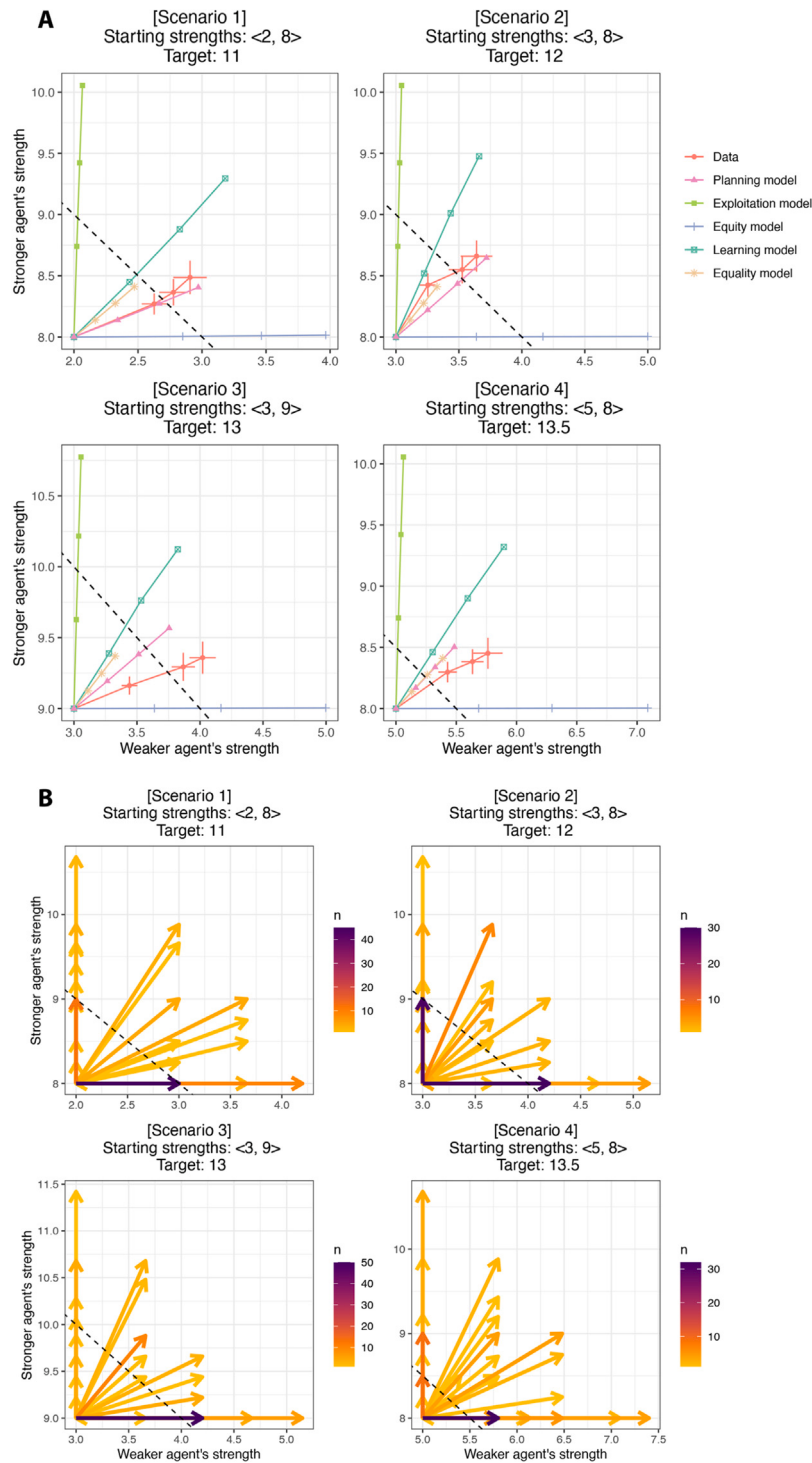


Fig. 2. Experiment 1 results. (A) Data and model predictions of agents' strength in each round of training. Each subplot corresponds to one scenario, where the x -axis shows the weaker agent's strength, the y -axis shows the stronger agent's strength, and each line traces agents' average strength in each training trial. The dashed black line marks the threshold where agents' combined strength exceeds the target box weight; intuitively, if participants are balancing the costs and benefits of training, they should train agents until their combined strength reaches this threshold, and no further. Error bars indicate 95% confidence intervals. (B) Individual participants' training trajectories. Each subplot denotes a scenario, and each arrow denotes a particular training strategy; the color of the arrow indicates how many participants adopted this strategy. For simplicity, here we represent each training strategy by showing only agents' strength before (origin) and after training (arrowhead). The dashed black line marks the threshold where agents' combined strength exceeds the target box weight.

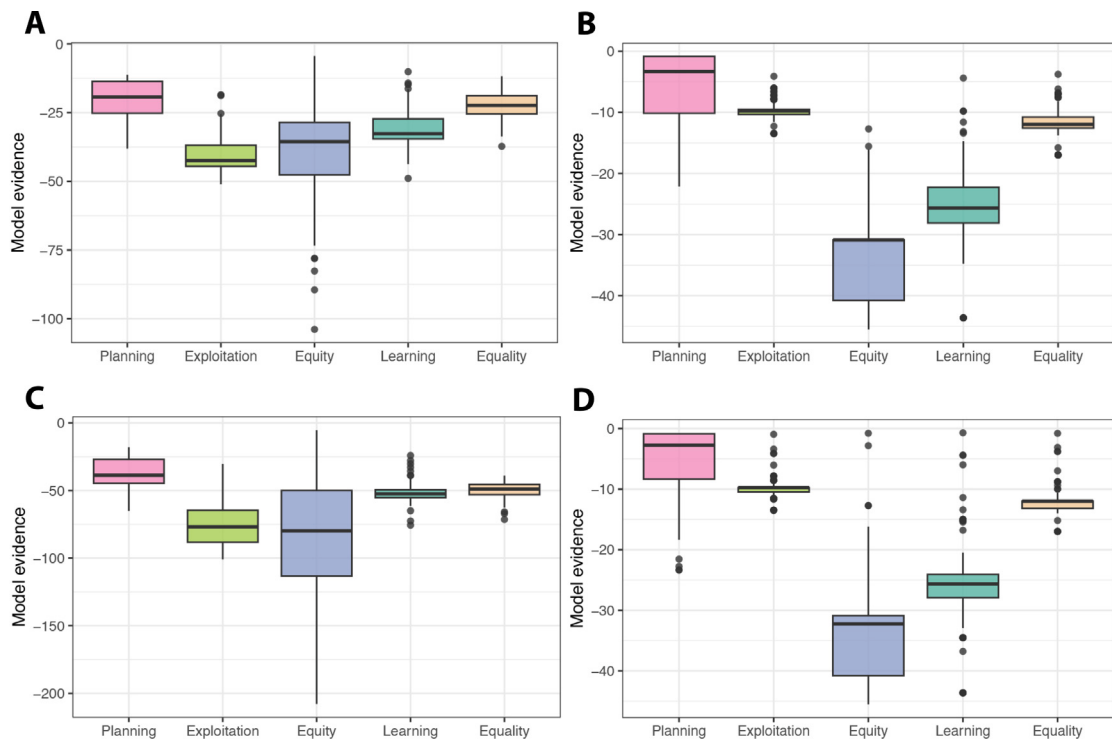


Fig. 3. Distribution of model evidences for each participant (approximated as $-0.5BIC$) in (A) Experiment 1, (B) Experiment 2, (C) Experiment 3, and (D) Experiment 4.

Finally, we took a closer look at individual participants' data (Fig. 2B). Each arrow corresponds to a particular training strategy, and the color of the arrow indicates how many participants adopted this strategy. For simplicity, here we represent each training strategy by showing only agents' strengths before (origin) and after training (arrowhead). Intuitively, diagonal arrows indicate that the participant trained both agents, horizontal arrows indicate training the weaker agent only, and vertical arrows indicate training the stronger agent only. While some participants trained agents more than what was necessary for goal achievement – either training both agents in turn or training a single agent over multiple rounds – overall, we found that the modal trajectories for all the scenarios show that participants trained either the stronger agent or the weaker agent until agents' combined strength reached the threshold for success (indicated by the dashed black line). Here, three out of four scenarios in Experiment 1 required only one round of training for success. In Experiment 3, we explored scenarios where three rounds of training are needed.

3.3. Discussion

In this experiment, we compared the planning model and four alternative models to participants' training strategies using random-effects Bayesian model selection. Overall, the planning model is the closest to the data both qualitatively and quantitatively and occurs the most frequently in the population.

4. Experiment 2

Experiment 1 addressed the question of how people decide which team members to train and how to train them. However, it does not address the problem of whom to recruit to the team in the first place. This question is important because, in real-world scenarios, a team does not necessarily need to recruit members who are already capable of a collaborative task to succeed at it; they can also recruit members who can be trained to be good enough for the task. For example, a company can either choose to poach a highly-skilled employee, or hire someone less experienced who will require training on the job. In Experiment 2, we modified our design to study how training considerations affect people's decisions about whom to recruit. Instead of giving participants two agents to train as in Experiment 1, here we asked participants to first hire two contestants from a pool of candidates with varying strengths and hiring costs that scale with strength and then train contestants to lift a heavy box together. We designed three scenarios with varying target box weights to test which model best predicts which candidates participants tend to hire. Intuitively, the planning model selects teams based on the target, forming different teams as needed for different scenarios, while alternative models each tend to select the same teams in all scenarios.

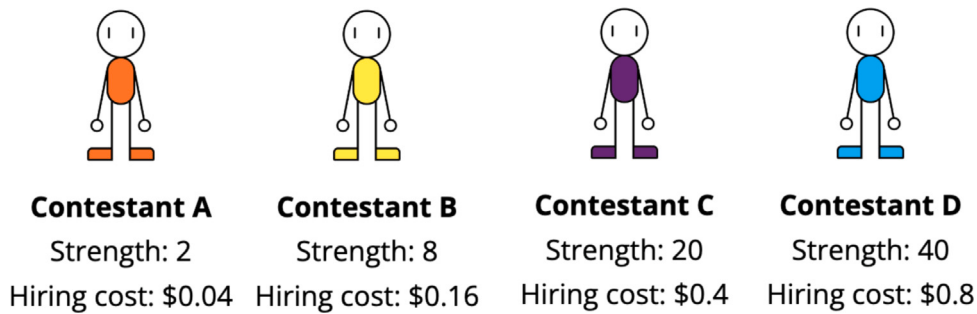


Fig. 4. The four candidates in Experiment 2. Hiring costs are proportional to agents' strength.

4.1. Materials and methods

4.1.1. Participants

We recruited 100 participants via Amazon's Mechanical Turk platform (MTurk). Participants received a base pay of \$1 and a potential bonus payment of up to \$6. The bonus payment depended on which contestants participants selected to train and complete the task, and whether the lifting turned out successful. Specifically, for every successful lift, participants received \$2; for every strength unit they used, \$0.02 was deducted from their bonus payment, regardless of the lift outcome. There were no costs associated with training boxes. The experiment was approved by the Harvard Institutional Review Board and pre-registered at https://aspredicted.org/6LJ_2D9.

4.1.2. Procedure

The game show consisted of three contests. At the start of each contest, participants were shown four agents who auditioned to appear in the game show. Each potential contestant had a different starting strength and cost a different amount of money to hire. Fig. 4 shows the contestants and the hiring fees; in general, stronger contestants cost more to hire. In each contest, participants chose two contestants and trained them separately to prepare them to lift a heavy box together. Once the participants selected two contestants, the training procedure was identical to Experiment 1.

4.2. Results

We compared the proportion of participants that chose each team to the model-predicted probabilities of choosing each team. As shown in the first row of Fig. 5A, participants hired agents who were either just strong enough for the task (e.g., agents with strengths 8 and 20 for a target of 24) or that could become strong enough for the task with some training (e.g., agents with strengths 2 and 20 for a target of 24). When both options were available, participants favored the latter – hiring agents that needed training over agents that did not – to reduce hiring costs. The planning model (second row) captured these response patterns. Similarly to participants, the planning model balanced the hiring costs against the reward of successful collaborations. Thus, it preferred teams that would succeed (either with or without training) over teams that would not, and it preferred teams that were cheaper to recruit among teams that would succeed. The alternative models (last four rows), on the other hand, did not predict these patterns in the data. Instead, the alternative models each preferred the same teams in all scenarios, regardless of whether this team composition was appropriate for the target box weight, and all alternative models had some tendency to hire weaker agents to lower hiring costs. This tendency was strongest in the learning model, which selected the two weakest agents to maximize strength gain, and weakest in the exploitation model, which had some tendency to hire the strongest agent to maximize the strongest agent's strength after training. By contrast, the equity model preferred to hire the single weakest agent to minimize the gap between the strongest and weakest agents, and the equality model dispreferred hiring the strongest agent due to hiring costs; however, neither of them had strong preferences among the remaining agents.

To evaluate each model's quantitative fit to the data, we computed correlations between the data and model predictions. The planning model has a very strong correlation with the data (Pearson's $r = 0.951, p < .0001$), whereas the alternative models all have weak correlation with the data: the exploitation model ($r = -0.256, p = .31$), equity model ($r = -0.207, p = .41$), learning model ($r = -0.195, p = .44$), and equality model ($r = -0.218, p = .38$).

Consistent with Experiment 1, we also compared each of the models to behavioral data using random-effects Bayesian model selection. We found that the planning model has a protected exceedance probability of approximately 1, suggesting that the planning model is the most likely model in the population. The planning model also had the highest model evidence (Fig. 3B).

We additionally plotted individual participants' training trajectories. For every scenario, we zoomed in on the training trajectories of the modal teams: (2, 20) in Scenario 1, (8, 40) in Scenario 2, and (20, 40) in Scenario 3 (Fig. 5B). As before, each arrow represents agents' pre-training and post-training strength based on a specific training strategy. Diagonal arrows indicate that the participant trained both agents, horizontal arrows indicate training the weaker agent only, and vertical arrows indicate training the stronger agent only. The color of each arrow shows the number of participants who adopted each strategy. We see that the majority of participants trained only one agent, to the point where agents' combined strength exceeded the threshold for success (the dashed black line).

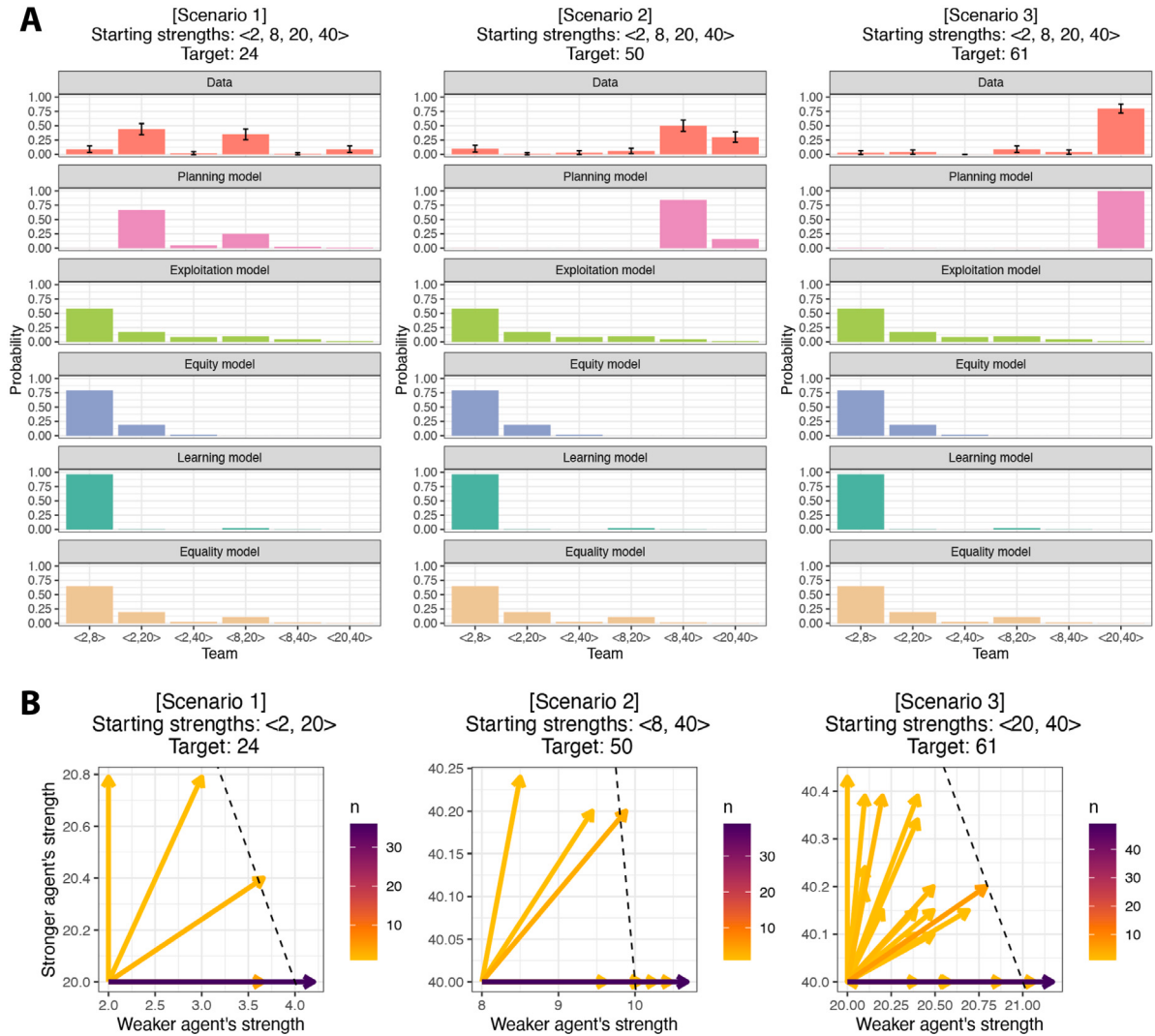


Fig. 5. Experiment 2 results. (A) Data and model predictions, showing the probability of choosing each team. Each subplot corresponds to one scenario; the brackets on the x-axis labels refer to the strength combinations of the two hired agents. Error bars indicate 95% confidence intervals of proportions. (B) Individual participants' training trajectories of the modal teams. Each subplot denotes a scenario, and each arrow denotes a particular training strategy; the color of the arrow indicates how many participants adopted this strategy. For simplicity, here we represent each training strategy by showing only agents' strength before (origin) and after training (arrowhead). The dashed black line marks the threshold where agents' combined strength exceeds the target box weight.

4.3. Discussion

In Experiment 2, we extended our task setup and models to the problem of team selection. The planning model captures participants' decisions about training *and* recruitment both qualitatively and quantitatively, whereas the alternative models deviate from the data substantively. Put together, these results suggest that people anticipate how their collaborators' competence will change and what problems they will face in the future when they decide whom to recruit.

5. Experiment 3

To test the generalizability of our findings, in Experiments 3 and 4, participants completed tasks that had the same underlying structure as Experiments 1 and 2, but were presented in a different context. Specifically, instead of asking participants to hire and train contestants of different strengths for a physical task, we asked them to recruit and train students with different math abilities for a math Olympiad. Though the underlying task structure remains the same, these two contexts differ from one another in potentially important ways; for example, rather than training contestants in one-off games, participants in this task assumed the role of an educator, in a context where planning may be superseded by concerns about allocating training equally or to those who

need it the most. These experiments thus provide a test of whether the cognitive processes that underlie decisions about recruitment and training apply more broadly to sequential social decision-making situations, beyond one context. Experiment 3 is a replication of Experiment 1 and Experiment 4 is a replication of Experiment 2. We included all combinations of starting strengths and target box weights used in the first two experiments; here, they were reframed as starting math abilities and target problem difficulties, respectively.

In Experiment 3, to discriminate between the planning model and equality model, we added four new scenarios (Scenarios 5–8) where more training is required to reach the target, as opposed to Scenarios 1–4 (same as Experiment 1) where only a little training is needed. The new scenarios had the same starting math abilities as the original scenarios, but higher target problem difficulties. In these scenarios, the planning model predicts that participants will train agents more as the target increases, whereas alternative models predict the same training strategies despite the target increase. If participants weigh the costs of training against the benefits of successful collaboration, then we should expect that they should also train agents differently to achieve different targets.

5.1. Materials and methods

5.1.1. Participants

We recruited 98 participants via Amazon's Mechanical Turk platform (MTurk). Participants received a base pay of \$2 and a potential bonus payment of up to \$8. The bonus payment depended on the problem sets participants used to train the student contestants and whether the student contestants won. Specifically, for every win, participants received \$1; for every difficulty level unit of the problem sets they used, \$0.04 was deducted from their bonus payment, regardless of the contest outcome. The experiment was approved by the Harvard Institutional Review Board and pre-registered at https://aspredicted.org/LJ7_YJD.

5.1.2. Procedure

Participants were instructed to imagine themselves as high school math coaches. Their school was sending eight different pairs of students to participate in eight different math Olympiads, and the participants' job was to train students to prepare them to work out a difficult math problem together to win the contest. Students can get better at math by completing problem sets; their abilities, indexed by their "math levels", increase by the level of effort they put into solving the problem sets, and a problem set that is too hard for them will not increase their math levels. The training procedure was the same as in Experiment 1.

5.2. Results

Fig. 6 visualizes the data and model predictions in a series of 2-D plots. Plots on the same row show scenarios with the same starting math levels but different targets. As in Experiment 1, each line shows how both agents' math levels changed, on average, over the course of training; the math level of the agent who was worse at math is plotted on the x-axis, and the math level of the agent who was better at math is plotted on the y-axis. Each point in the participants' responses denotes the participant-averaged math level of the two agents in every round of training. Each point in the model predictions denotes the weighted average math level calculated from Eq. (9) using γ fit to individual participants. The dashed black lines indicate the minimum total math level required to win the contest. Participants' average responses aligned closely with the planning model; both have endpoints just past the dashed line, indicating that participants tended to train agents to be just good enough for the task, to maximize rewards and minimize training costs. By contrast, the alternative models failed to capture this qualitative pattern: The exploitation model trained the agent who was better at math as much as possible, the equity model trained the agent who was worse at math as much as possible, and the learning model trained the two agents for the largest possible gain in math levels. The equality model, on the other hand, did not train agents to be good enough to win the contest in seven out of eight scenarios. In addition, by horizontally comparing scenarios with the same starting math levels but different targets, we see that participants adjusted their training decisions when the target changed. The planning model was the only one that captured this pattern; all alternative models predicted the same training strategies when the target changed.

We computed correlations between the participant-averaged data and model predictions of agents' math level at every step of training. The planning model has the highest correlation with the data (Pearson's $r = 0.999, p < .0001$), followed by the equality model ($r = 0.998, p < .0001$), learning model ($r = 0.996, p < .0001$), exploitation model ($r = 0.992, p < .0001$), and equity model ($r = 0.983, p < .0001$). Note that all correlations between model predictions and data are high because agents' math levels increase monotonically, which obscures differences in how well each model captures qualitative patterns in the data. Therefore, we conducted a formal model comparison using random-effects Bayesian model selection to quantitatively assess how well each model captures participants' training strategies. We found that the planning model's protected exceedance probability was approximately 1, meaning that the planning model is the most frequently occurring model in the population. In individual participants, the planning model also had the highest model evidence (Fig. 3C). Compared to Experiment 1, here the differences between the evidence for the planning model and equality model were more pronounced because of the four new scenarios designed to discriminate between them.

Finally, looking at individual participants' training trajectories (Fig. 7), we see that, overall, the modal trajectories in all the scenarios were to train agents to be just good enough to win. More specifically, in scenarios where only a little training was required for winning (Scenarios 1–4), participants tended to train just one agent, whereas when a lot of training was needed (Scenarios 5–8), participants tended to either solely train the agent who was worse at math, or to provide training to both agents.

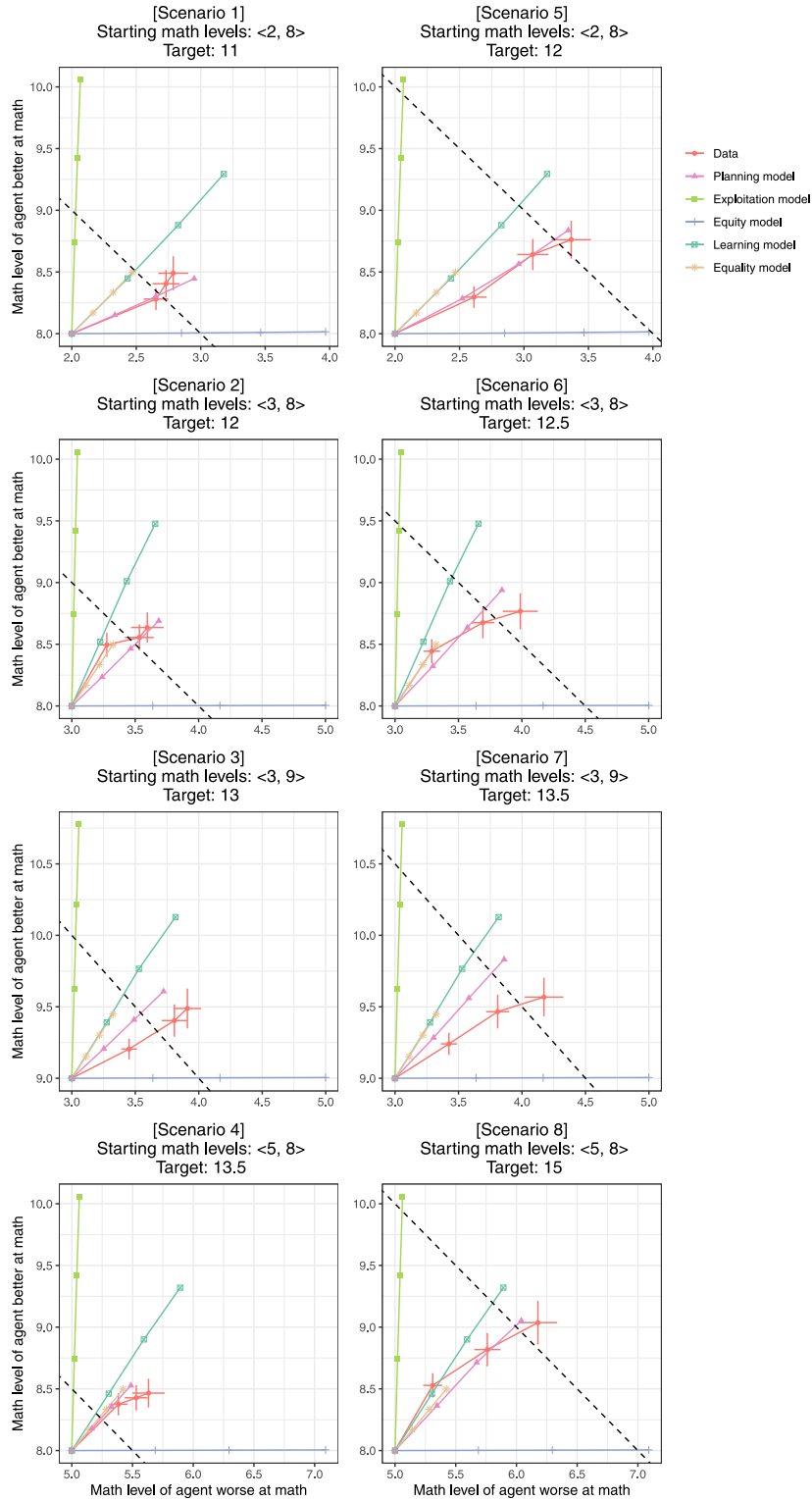


Fig. 6. Experiment 3 data and model predictions of agents' math level in every round of training. Each subplot corresponds to one scenario, where the x-axis shows the math level of the agent who was worse at math, the y-axis shows the math level of the agent who was better at math, and each line traces agents' average math level in each training trial. The dashed black line marks the threshold where agents' combined math level exceeds the target difficulty level of the contest; intuitively, if participants are balancing the costs and benefits of training, they should train agents until their combined math level reaches this threshold, and no further. Error bars indicate 95% confidence intervals.

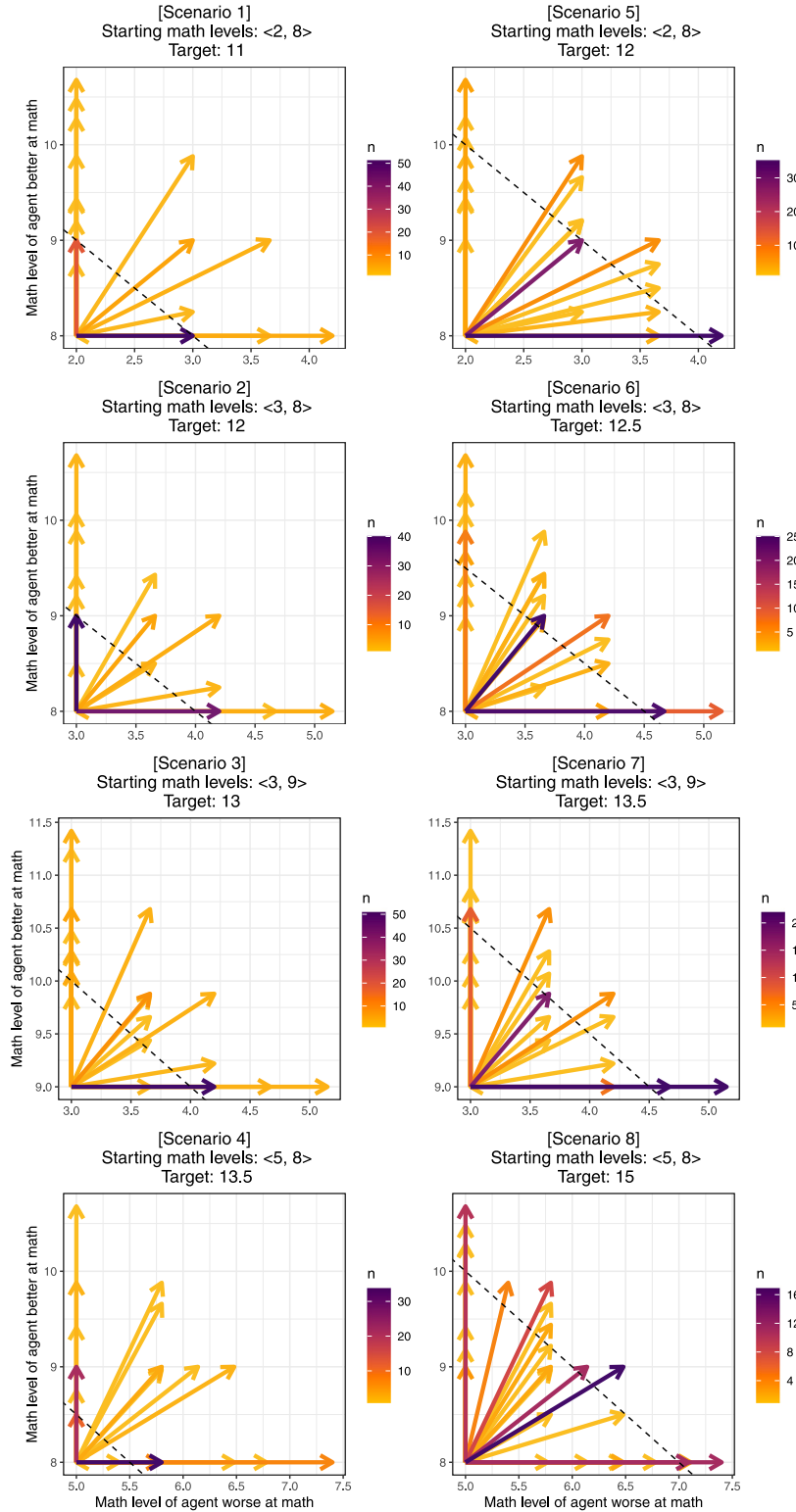


Fig. 7. Experiment 3 individual participants' training trajectories. Each subplot denotes a scenario, and each arrow denotes a particular training strategy; the color of the arrow indicates how many participants adopted this strategy. For simplicity, here we represent each training strategy by showing only agents' math levels before (origin) and after training (arrowhead). The dashed black line marks the threshold where agents' combined math level exceeds the target difficulty level of the contest.

5.3. Discussion

In Experiment 3, we replicated the findings in Experiment 1 and showed that the findings generalized to a new and more education-orientated framing—training students for a math Olympiad. Each alternative model captures aspects of participants' responses – e.g., the equality model mirrors participants' responses in contexts where only a little training is required – but only the planning model captures the full pattern of the data.

6. Experiment 4

Experiment 3 successfully replicated the results in Experiment 1 in a new context. In Experiment 4, we replicated Experiment 2 to ask whether the planning model also captures participants' decisions about whom to recruit for math Olympiads.

6.1. Materials and methods

6.1.1. Participants

We recruited 99 participants via Amazon's Mechanical Turk platform (MTurk). Participants received a base pay of \$1 and a potential bonus payment of up to \$6. The bonus payment depended on which students participants recruited, and whether training enabled them to win the contest. Specifically, for every win, participants received \$2; for every math level unit in the students they recruited, \$0.02 was deducted from their bonus payment, regardless of the contest outcome. There were no costs associated with problem sets. The experiment was approved by the Harvard Institutional Review Board and pre-registered at https://aspredicted.org/T6Z_XCD.

6.1.2. Procedure

As in Experiment 3, participants imagined themselves as high school math coaches. Their school was planning to send pairs of students to participate in different math Olympiad contests; participants were asked to first recruit two students for each contest among four students who auditioned to participate, and then train them to prepare them to work out a difficult math problem together. Each student had a different starting math level and required a different amount of money to recruit. The stimuli and structure of the game were the same as in Experiment 2.

6.2. Results

As in Experiment 2, overall, participants recruited students who were either just good enough for the contest, or could win the contest with the least training. Among the two options, participants favored recruiting students who were not good enough and training them to win (see the first row of Fig. 8A), in order to minimize hiring costs. The planning model was the only model that captured these patterns (second row); none of the alternative models based their choices on whether a team would win the contest or not (third to last rows).

Correlation analyses revealed that the planning model has a very strong correlation with the data (Pearson's $r = 0.974, p < .0001$), whereas the alternative models all have weak correlation with the data (exploitation model: $r = -0.289, p = .25$; equity model: $r = -0.246, p = .32$; learning model: $r = -0.234, p = .35$; equality model: $r = -0.257, p = .30$). Random-effects Bayesian model selection further showed that the planning model has a protected exceedance probability of approximately 1, suggesting that the planning model is the most likely model in the population. The planning model also had the highest model evidence (see Fig. 3D).

Zooming in on individual participants' training trajectories of the modal teams ($\langle 2, 20 \rangle$ in Scenario 1, $\langle 8, 40 \rangle$ in Scenario 2, and $\langle 20, 40 \rangle$ in Scenario 3 (Fig. 8B), we see that the majority of participants trained the agent who was worse at math exclusively, and just enough to win the contest.

6.3. Discussion

Experiment 4 replicated Experiment 2 in a new context where participants recruited and trained students for a math Olympiad. As in Experiment 2, the planning model captures participants' decisions about training *and* recruitment both qualitatively and quantitatively, but none of the alternative models does. These results suggest that the cognitive processes we uncover apply more broadly to sequential social decision-making situations, beyond the context of hiring and training workers.

7. General discussion

Members of a collaboration often gain and hone their skills over the course of collaborative tasks—for example, a student may become a stronger writer while drafting a manuscript, or an employee may learn to carry out her tasks more effectively on the job. However, little is known about how people account for this plasticity in collaborative decision-making. In the current work, we conducted four experiments to investigate how people decide which team members to train and how to train them (Experiments 1 and 3), and whom to recruit to a team (Experiments 2 and 4). We compared a planning model that anticipates how agents'

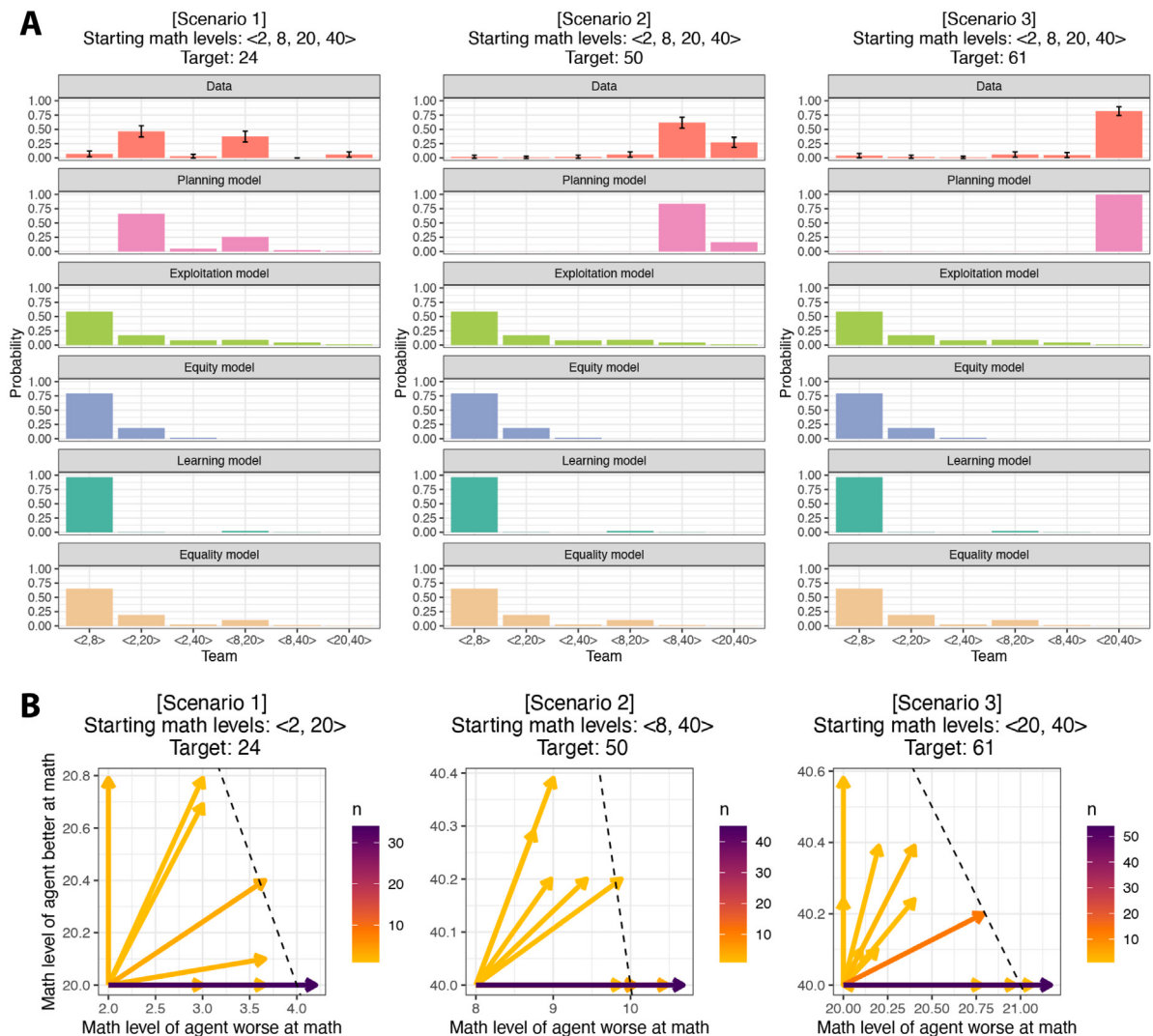


Fig. 8. Experiment 4 results. (A) Data and model predictions, showing the probability of choosing each team. Each subplot corresponds to one scenario; the brackets on the x-axis labels refer to the math levels of the two recruited students. Error bars indicate 95% confidence intervals of proportions. (B) Individual participants' training trajectories of the modal teams. Each subplot denotes a scenario, and each arrow denotes a particular training strategy; the color of the arrow indicates how many participants adopted this strategy. For simplicity, here we represent each training strategy by showing only agents' math levels before (origin) and after training (arrowhead). The dashed black line marks the threshold where agents' combined math level exceeds the target difficulty level of the contest.

competence will change through training to four alternative models that do not plan ahead: an exploitation model, an equity model, a learning model, and an equality model. By comparing human judgments to model predictions, we found that the planning model is the only model that captures the qualitative patterns in participants' responses and provides the best quantitative fit overall. These findings were consistent in both a physical task where contestants lifted a box together and an educational task where students solved a math problem together, demonstrating that the cognitive processes we uncover apply broadly to sequential social decision-making situations. Our results suggest that, in collaborations, people do not solely recruit the most competent agents for the job, nor do they invest solely in training while losing sight of the goal; instead, people plan ways to optimize training in the service of collaboration, maximizing the benefits to collaborators while minimizing the costs of training.

Though the alternative models performed worse than the planning model in our study, there are scenarios where the planning model makes predictions similar to the alternative models, and the heuristics of these alternative models can serve as valid shortcuts to the decision problem. For example, when the time horizon is very short, people might simply pick the more competent agents since there is less payoff for skill development through training. In the face of a difficult task, people might opt for maximizing total improvement. And when a task requires that each team member execute an individual subtask, people might prioritize training the less competent agents. These heuristics may also be a reasonable way of hedging one's bets when it is difficult to anticipate exactly to what degree a person's competence will improve with training, or the exact difficulty of the problems that they will face in the

future. However, even if these heuristics may explain how people make decisions about training within particular contexts, they do not explain why specific heuristics are appropriate within specific scenarios. By contrast, the planning model captures variability in participants' decisions across different training contexts. Thus, it is possible that the planning model provides a theoretical framework to understand the *variability* in training strategies across different, idealized scenarios. Future work needs to be done to understand how these planning models fare under uncertainty, and to develop process models that approximate planning over longer-running collaborations with larger action spaces.

Understanding variability in training strategies sheds light on the relationship between teaching and collaboration. Teaching is often defined as a modification in a teacher's behavior that enables a learner to acquire a skill or knowledge more rapidly than they would on their own, at a cost or for no immediate benefit to the teacher (Caro & Hauser, 1992). While this definition helps to distinguish teaching from instrumental behavior, what it misses is that the two are often intertwined: Humans often pass on skills through collaborative relationships that also benefit the mentor, such as guilds, apprenticeships, and job training. The current work therefore motivates new questions about how teachers and learners who are enmeshed in collaborative relationships balance the pedagogical benefits to a learner against the constraints and costs of the collaborative task. Viewed in the context of this work, by assigning tasks to the most competent agent, the exploitation model prioritizes meeting the demands of the collaborative task without considering who needs training, whereas the equity model and the learning model focus solely on the benefits of training and ignores the demands of the task. Moving forward, how the nature of the collaborative relationship – such as the size of the collaborative group, or the time horizon of the collaboration – affect the degree to which we invest in training collaborators, and how these additional, instrumental concerns affect decisions about who and what is most helpful to teach, remain open questions for future research.

Our framework used a simple setting to study questions related to training and competence. However, it can be straightforwardly extended to address more complex settings. In our setup, team members had only one task to complete, whereas, in multi-task collaborative settings, team members typically have to complete a variety of tasks. For example, a research group often works on multiple different projects simultaneously and, within each project, there are usually a myriad of tasks to complete, such as devising theories, running experiments, analyzing data, and writing up the findings. Therefore, good teams need to possess some degree of versatility, which may either be from each individual being able to fit multiple roles or from building up a diverse team where different people can satisfy the demands of different tasks. Besides being versatile, realistic settings also require that teams be robust to turnover. Our design did not include attrition, but in reality, team members might drop out at any time—for example, a researcher might leave the group, and someone else on the team might need to pick up where she left off. It is thus important to train people with skills that are generally adaptive, rather than just the ability to do one thing. Further work is needed to understand how teams figure out the optimal degree of specialization.

Another limitation is that we made improvements in competence directly observable and quantifiable. These assumptions serve their purpose in our study, but realistic settings are more complex. For example, completing a project will not make a researcher look evidently more knowledgeable right away, and the researcher herself might find it difficult to tell exactly how many research skills she gained over the past year. In addition, we simplified how training changes competence by only considering increases in competence and assuming that competence increases by the level of effort people put forth in training. In the real world, people's performance may also *decrease* as a function of fatigue and boredom, and the relationship between effort and subsequent improvements is more complex and can vary from skill to skill.

An open question for future work is how well these findings scale in more realistic conditions. One intriguing direction is to incorporate collaborators' theory of mind into the models (e.g., Xiang et al., 2023b). We deliberately simplified the theory of mind requirements in our task so that we could focus on understanding how people use their intuitive theory of mind to guide training decisions. However, it leaves open the question of whether and how participants take into consideration their collaborators' beliefs and desires when making hiring and training decisions. For example, people differ in their willingness to exert effort, and if we allow that to vary in Experiment 3, we might expect students of the same math level to benefit from the same problem set differently depending on how diligent they are. Another possible extension of this work is to apply the same task structure to hiring and training real people. Our task used artificial agents; this allowed us to study the plasticity of competence in a tightly controlled experimental context, where changes in competence happen in a matter of seconds right after an action, are directly visible, and can be accurately anticipated. Even though we made an effort to simulate the experience of collaborating with real people (e.g., by framing agents as contestants in a game show or students in a math Olympiad), and our previous work (Xiang et al., 2023a, 2023b) and many studies in the field (e.g., Baker et al., 2017; Gerstenberg et al., 2018; Jara-Ettinger et al., 2016; Liu et al., 2017; Ullman et al., 2009; Vélez & Gweon, 2020) have demonstrated that adults, children, and even infants are capable of viewing artificial agents as humans, it is possible that participants might weigh their own payoffs and other factors differently, such as caring more about collaborators' learning, or about bridging the competence gap between them. Because improving competence takes time, multiple rounds of evident improvement might not be feasible in lab experiments; researchers may need to resort to large-scale long-term observational studies, such as tracking the training decisions of a company or a sports team.

Here, we present a planning model of collaboration that makes decisions about whom to recruit and how to train them by anticipating how collaborators' competence will change and what problems they will face in the future. This work provides a first step towards understanding how humans facilitate and invest in other people's learning in a wide range of collaborative relationships, e.g., as pupils, apprentices, and trainees.

CRediT authorship contribution statement

Yang Xiang: Writing – review & editing, Writing – original draft, Visualization, Project administration, Methodology, Investigation, Funding acquisition, Formal analysis, Data curation, Conceptualization. **Natalia Vélez:** Writing – review & editing, Writing – original draft, Validation, Supervision, Methodology, Conceptualization. **Samuel J. Gershman:** Writing – review & editing, Writing – original draft, Validation, Supervision, Methodology, Conceptualization.

Data availability

We have made our data and code publicly available at https://github.com/yyxiang/competence_training.

Acknowledgments

We are grateful to Gudela Grote for helpful discussions. This research was funded by a Smith Fund grant to Y.X. from the Harvard University Department of Psychology, United States of America and an NIMH, United States of America K00 award (award K00MH125856) to N.V.

References

- Baer, C., & Odic, D. (2022). Mini managers: Children strategically divide cognitive labor among collaborators, but with a self-serving bias. *Child Development*, 93(2), 437–450.
- Baker, C. L., Jara-Ettinger, J., Saxe, R., & Tenenbaum, J. B. (2017). Rational quantitative attribution of beliefs, desires and percepts in human mentalizing. *Nature Human Behaviour*, 1(4), 0064.
- Baker, C. L., Saxe, R., & Tenenbaum, J. B. (2009). Action understanding as inverse planning. *Cognition*, 113(3), 329–349.
- Bartel, A. P. (1994). Productivity gains from the implementation of employee training programs. *Industrial Relations: A Journal of Economy and Society*, 33(4), 411–425.
- Bentham, J. (1970). In J. H. Burns, & H. L. A. Hart (Eds.), *An Introduction to the Principles of Morals and Legislation* (1789). London.
- Bishop, C. M. (2006). *Pattern recognition and machine learning*. Springer.
- Caro, T. M., & Hauser, M. D. (1992). Is there teaching in nonhuman animals? *The Quarterly Review of Biology*, 67(2), 151–174.
- Curioni, A. (2022). What makes us act together? On the cognitive models supporting humans' decisions for joint action. *Frontiers in Integrative Neuroscience*, 16, Article 900527.
- Daw, N. D., & Dayan, P. (2014). The algorithmic anatomy of model-based evaluation. *Philosophical Transactions of the Royal Society, Series B (Biological Sciences)*, 369, Article 20130478.
- Dhami, M. K. (2003). Psychological models of professional decision making. *Psychological science*, 14(2), 175–180.
- Elnaga, A., & Imran, A. (2013). The effect of training on employee performance. *European Journal of Business and Management*, 5(4), 137–147.
- Gerstenberg, T., Ullman, T. D., Nagel, J., Kleiman-Weiner, M., Lagnado, D. A., & Tenenbaum, J. B. (2018). Lucky or clever? From expectations to responsibility judgments. *Cognition*, 177, 122–141.
- Hertwig, R., Davis, J. N., & Sulloway, F. J. (2002). Parental investment: How an equity motive can produce inequality. *Psychological bulletin*, 128(5), 728.
- Ho, M. K., Saxe, R., & Cushman, F. (2022). Planning with theory of mind. *Trends in Cognitive Sciences*, 26(11), 959–971.
- Huys, Q. J., Lally, N., Faulkner, P., Eshel, N., Seifritz, E., Gershman, S. J., Dayan, P., & Roiser, J. P. (2015). Interplay of approximate planning strategies. *Proceedings of the National Academy of Sciences*, 112(10), 3098–3103.
- Jara-Ettinger, J., Gweon, H., Schulz, L. E., & Tenenbaum, J. B. (2016). The naïve utility calculus: Computational principles underlying commonsense psychology. *Trends in Cognitive Sciences*, 20(8), 589–604.
- Liu, S., Ullman, T. D., Tenenbaum, J. B., & Spelke, E. S. (2017). Ten-month-old infants infer the value of goals from the costs of actions. *Science*, 358(6366), 1038–1041.
- Magid, R. W., DePascale, M., & Schulz, L. E. (2018). Four-and 5-year-olds infer differences in relative ability and appropriately allocate roles to achieve cooperative, competitive, and prosocial goals. *Open Mind*, 2(2), 72–85.
- Messick, D. M. (1993). *Equality as a decision heuristic*. Cambridge University Press.
- Mill, J. S. (2016). Utilitarianism. In *Seven masterpieces of philosophy* (pp. 329–375). Routledge.
- Morton, R. W., Colenso-Semple, L., & Phillips, S. M. (2019). Training for strength and hypertrophy: An evidence-based approach. *Current Opinion in Physiology*, 10, 90–95.
- Nda, M. M., & Fard, R. Y. (2013). The impact of employee training and development on employee productivity. *Global Journal of Commerce and Management Perspective*, 2(6), 91–93.
- Payne, J. W., Bettman, J. R., & Johnson, E. J. (1993). *The adaptive decision maker*. Cambridge University Press.
- Payne, J. W., Bettman, J. R., & Luce, M. F. (1996). When time is money: Decision behavior under opportunity-cost time pressure. *Organizational Behavior and Human Decision Processes*, 66(2), 131–152.
- Premack, D., & Woodruff, G. (1978). Does the chimpanzee have a theory of mind? *Behavioral and Brain Sciences*, 1(4), 515–526.
- Rawls, J. (2001). *Justice as fairness: A restatement*. Harvard University Press.
- Rigoux, L., Stephan, K. E., Friston, K. J., & Daunizeau, J. (2014). Bayesian model selection for group studies—revisited. *Neuroimage*, 84, 971–985.
- Smith, A. (1937). *The wealth of nations [1776]*. New York: Modern Library.
- Török, G., Pomiechowska, B., Csibra, G., & Sebanz, N. (2019). Rationality in joint action: Maximizing efficiency in coordination. *Psychological Science*, 30(6), 930–941.
- Török, G., Stanciu, O., Sebanz, N., & Csibra, G. (2021). Computing joint action costs: Co-actors minimize the aggregate individual costs in an action sequence. *Open Mind*, 5, 100–112.
- Ullman, T., Baker, C., Macindoe, O., Evans, O., Goodman, N., & Tenenbaum, J. (2009). Help or hinder: Bayesian models of social goal inference. In *Advances in neural information processing systems: vol. 22*.
- Vélez, N., & Gweon, H. (2020). Preschoolers use minimal statistical information about social groups to infer the preferences and group membership of individuals. In *CogSci*.
- Vélez, N., & Gweon, H. (2021). Learning from other minds: An optimistic critique of reinforcement learning models of social learning. *Current Opinion in Behavioral Sciences*, 38, 110–115.

- Warneken, F., Steinwender, J., Hamann, K., & Tomasello, M. (2014). Young children's planning in a collaborative problem-solving task. *Cognitive Development*, 31, 48–58.
- Xiang, Y., Landy, J., Cushman, F. A., Vélez, N., & Gershman, S. J. (2023a). Actual and counterfactual effort contribute to responsibility attributions in collaborative tasks. *Cognition*, 241, Article 105609.
- Xiang, Y., Vélez, N., & Gershman, S. J. (2023b). Collaborative decision making is grounded in representations of other people's competence and effort. *Journal of Experimental Psychology: General*, 152(6), 1565–1579.