

Cognitive Science 50 (2026) e70247
© 2026 Cognitive Science Society LLC.
ISSN: 1551-6709 online
DOI: 10.1111/cogs.70247

Training Collaborators for Effective Division of Labor

Fangze Tian,^a Samuel J. Gershman,^{a,b} Yang Xiang^a

^a*Department of Psychology, Harvard University*

^b*Center for Brain Science, Harvard University*

Received 30 January 2026; received in revised form 21 June 2026; accepted 1 July 2026

Abstract

By dividing labor, collaboration brings together the complementary strengths of individuals. Yet, division of labor benefits collaboration not only by assigning people to tasks that match their competences, but by allowing them to improve “on-the-job” through training and practice. Thus, optimal division of labor requires anticipating how each individual will develop and what roles they will eventually be best suited for. Existing research on collaboration, however, rarely considers this prospective dimension. To address this gap, we studied how humans make training decisions, while manipulating the long-term consequences of those decisions. Across three experiments ($N = 600$), participants trained two military defense teams to counter two types of attacks (land and air), where the long-term goals and the teams’ relative competences varied. Participants made a sequence of training decisions before assigning teams to roles in the final battle. Overall, participants divided labor and trained collaborators by considering how each team’s competence would develop, and whether they would eventually meet task demands. These patterns were best captured by a Planning model that trained collaborators based on the expected outcome at the time of deployment. The Planning model outperformed several heuristic alternatives that optimized training based on current competence, learning potential, fairness, or versatility. Together, these findings provide a first step toward understanding how training unlocks one of the central benefits of dividing labor—complementary competences. People do not simply match existing competences to tasks; rather, they actively cultivate individual competences to support future specialization.

Keywords: Social cognition; Computational modeling; Collaboration; Division of labor; Theory of Mind

Correspondence should be sent to Yang Xiang, Department of Psychology, Harvard University, William James Hall, 33 Kirkland St., Cambridge, MA 02138, USA. E-mail: yyx@g.harvard.edu

All rights reserved, including rights for text and data mining and training of artificial intelligence technologies or similar technologies.

1. Introduction

Division of labor is a fundamental principle of human collaboration. By distributing work across individuals with complementary competences, each person can contribute where they are most effective. Classical economic theory recognizes three key advantages of dividing labor, the first and foremost being that it allows individuals to improve competences through practice¹ (Smith, 1937; Yang & Ng, 1998). In *The Wealth of Nations*, Smith (1937) illustrated this advantage with an example of nail-making: a teenager who makes nails all day can easily outperform a smith who only occasionally forges them. Smith (1937) further argued that differences in competence across professions are not innate, but rather acquired as a result of repeated practice within specialized roles. The power of division of labor, therefore, lies not only in efficiently matching people to tasks, but in how such specialization allows individuals to become more competent over time and enhance their future contributions to the collaborative enterprise.

Extensive research on division of labor has examined how people decide who should do what in joint tasks. Work in economics and organizational science shows that efficient labor division depends on team composition and coordination structure (Cooper, Frost, Liu, & West, 2021; Haeussler & Sauermann, 2016, 2020; Powell, 2000). Studies in cognitive science highlight how members of a team mutually adapt to complement what they expect others to do (Curioni, 2022; Goldstone, Andrade-Lotero, Hawkins, & Roberts, 2024; Wahn & Schmitz, 2023; Xiang, Vélez, & Gershman, 2023), often sacrificing individual efficiency to maximize co-efficiency (Török, Pomiechowska, Csibra, & Sebanz, 2019; Török, Stanciu, Sebanz, & Csibra, 2021). Together, these lines of work reveal that people flexibly divide labor according to collaborators' competences, task demands, and external constraints. However, most studies explain division of labor as a problem of matching *static* competences to team goals. In reality, competences evolve through experience: as individuals practice particular roles, they improve their competence in those roles. From a computational perspective, collaboration involves not only optimizing the present but also anticipating how current decisions—such as whom to train and how to train them—will change future competences (Liu, Xiang, & Gershman, forthcoming; Xiang & Gershman, 2025) and, in turn, the effectiveness of future divisions of labor.

Training can enhance future performance, but the best way to train depends on anticipating how collaborators' competences will change with training and how those changes will support the anticipated division of labor. Recent computational work began to address this question through a collaborative box-lifting task, showing that people plan ahead and train collaborators to be just good enough for the collaborative needs (Xiang, Vélez, & Gershman, 2024). While this work provides a first step toward understanding how people optimize training in the service of collaboration, the study only considered a single dimension of competence—physical strength—and thus does not address the division of labor. With only one kind of competence, there is no meaningful specialization or roles. In contrast, real-world collaborations typically rely on multiple, distinct competence dimensions, where effective collaboration requires anticipating not only how much each partner should improve, but also in which dimension improvement will matter most for the collaboration. To understand how

training supports division of labor, we need a framework that addresses the more general and ecologically realistic problem of optimizing competence for the division of labor in multi-dimensional competence spaces.

In this paper, we examine how participants train and divide labor based on multi-dimensional competences. Across three experiments ($N = 600$) with varying task demands, participants make a series of training decisions before assigning agents to different roles in the final task. We develop and compare six computational models to characterize participants' behavior. The first is a forward-looking Planning model inspired by model-based approaches to sequential decision problems (Daw & Dayan, 2014; Huys et al., 2015). The Planning model trains agents to optimize long-term collaborative goal achievement according to the task demands. Extensive planning is computationally demanding, which may lead people to simplify their decisions using myopic heuristics (Huys et al., 2012; Meder, Nelson, Jones, & Ruggeri, 2019; Rieskamp et al., 1999). Accordingly, we test five models based on simple heuristics: (1) the Exploitation model prioritizes training the most competent agent; (2) the Learning model trains to maximize total competence increase; (3) the Equity model trains the weaker agent to minimize inequality between agents; (4) the Equality model is indifferent across training options; and (5) the General model trains agents to increase versatility by reducing the imbalance across an agent's different dimensions of competences. Overall, we found that the Planning model best captures participants' training decisions, suggesting that people train collaborators to optimize long-term collaborative goal achievement by taking into account the roles collaborators will eventually need to play, and how training supports the role expectations.

2. Method

All data, materials, analysis scripts, links to the online experiments, and preregistration are available at https://github.com/yyxiang/multi_dim_training. All experiments were preregistered at <https://aspredicted.org/w5hp37.pdf>.

2.1. Experiments

In the experiments, participants oversaw two military teams (Red and Blue), each with competence in two competence dimensions (land and air defense). As shown in Fig. 1, participants first trained the teams to improve their competences (training phase), then assigned them to defend against land and air attacks (execution phase). Across experiments, we varied the structure and demand of the task:

- In the dimension-based structure (Experiment 1), the goal was to maximize total defense. The attacks did not occur on the same day, so either team could handle either dimension.
- In the agent-based structure (Experiment 2), attacks occurred simultaneously, so each team could only cover one dimension.

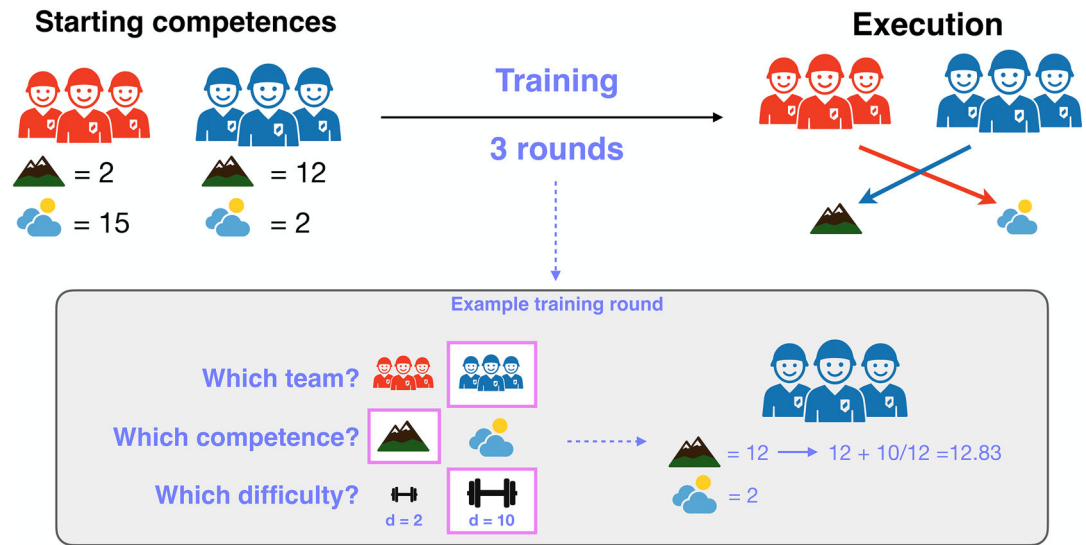


Fig. 1. Example scenario. Participants first observe the starting competences of each team in land and air defense, then train them over three rounds before assigning the teams to defend land and air attacks. Teams can improve their competence by exerting effort on training tasks.

- In the weakest-link structure (Experiment 3), attacks also occurred simultaneously, but instead of maximizing total defense, the goal was to minimize the loss—thus, maximizing the lower of the two assigned competences.

Within each experiment, we manipulated the teams' relative starting competences to examine how initial competence differences affect training and role assignment.

2.1.1. Participants

We recruited 200 participants per experiment ($N = 600$ in total) via Amazon Mechanical Turk using CloudResearch. Before beginning the main task, participants answered a few comprehension questions to ensure understanding of the task instructions and were allowed to proceed only after answering all questions correctly. We excluded participants who failed the attention check administered during the task, resulting in a final sample of $N = 589$ for analysis. All participants provided informed consent and were unaware of the study's hypotheses. The study took approximately 20 min to complete, and participants received a base payment of \$3.00 and a performance-based bonus of up to \$3.00. The study protocol was approved by the Harvard University Institutional Review Board.

2.1.2. Procedure

To ground the experimental setup under a realistic setting, we framed the task as a military defense scenario. Participants were told that they would oversee two teams of soldiers—the Red team and the Blue team—to defend their country against two types of attack—land

and air. This cover story creates a high-stakes context, encouraging participants to reason strategically about training and task assignment decisions.

At the start of the experiment, participants received detailed instructions, specifically on the structure of the task, the logic of the training and execution phases, and how performance would be evaluated. We instructed participants that the teams' competence increased by the amount of effort they exerted. We demonstrated the effects of these values through practice trials, where each participant observed three examples to learn how agents with different competences responded to different task difficulties. Example 1 showed that a team with competence 10 given a difficulty 10 task would exert 100% of effort, thus gain 1 in competence. Example 2 showed that this same team given a difficulty 2 task would exert only 20% effort, thus gain 0.2 in competence. Example 3 showed that a team with competence 6 given a difficulty 10 task would give up (i.e., exert 0% effort), and their competence would not change.

Trial structure: Each trial consisted of two phases: a *training* phase and an *execution* phase. During the training phase, participants were asked to make three sequential rounds of training decisions: in each round, they chose which team to train (Red or Blue), which competence dimension to train (Land or Air), and which training difficulty to use (2 or 10). In other words, they allocated training to one of four possible options: Red-Land, Red-Air, Blue-Land, or Blue-Air, and a training difficulty of either 2 or 10. All competence values were displayed numerically at the beginning of the trial, and updated after each training round to reflect the effects from training. Competence and training difficulty were expressed in the same units for intuitive interpretation.

After the three rounds of training, participants proceeded to the *execution* phase. In this phase, they were asked which team to assign for land defense (Red-land or Blue-land), and which for air defense (Red-air or Blue-air). These two questions were independent of each other, meaning that participants could choose to assign the same team for both defenses. Final task performance (i.e., total defense score) was computed according to the task structure described below, and the total defense score was displayed after each trial.

Experiment structure: The study manipulated two factors: *task structure* and *starting competence setup*. Task structure varied in the demand of the task and how final performance was computed; starting competence setups varied in the relative starting competences of the two teams. Within each competence setup, three *scenarios* varied the exact competence values, while maintaining the relative pattern. Thus, the design followed a 3 (task structure) $\times$ 3 (competence setup) $\times$ 3 (scenario) design. Participants were randomly assigned to one task structure (Experiments 1–3) and completed all nine trials. The order of trials was randomized, and the assignment of team colors (Red and Blue) was randomly counterbalanced across participants. Table 1 shows a complete list of the stimuli.

Each experiment had a distinct task structure that determined the collaborative goal. In the *dimension-based* structure (Experiment 1), the two attacks occurred at different times, allowing either team to defend either attack. In other words, one team can perform both defenses. Performance was determined by the total defense competence (the sum of the two selected competences): a higher sum meant a better score. In the *agent-based* structure (Experiment

Table 1
Starting competences for each starting competence setup × scenario trial

	Red-Land	Blue-Land	Red-Air	Blue-Air
Starting Competence Setup 1	<i>Red is more competent than Blue in both Land and Air.</i>			
Scenario 1	10	2	15	5
Scenario 2	15	2	20	5
Scenario 3	15	5	20	6
Starting Competence Setup 2	<i>Both teams are equal in Land, and Red is more competent in Air.</i>			
Scenario 1	2	2	10	2
Scenario 2	2	2	15	2
Scenario 3	5	5	15	5
Starting Competence Setup 3	<i>Red is more competent in Land, while Blue is more competent in Air.</i>			
Scenario 1	10	2	5	15
Scenario 2	5	2	5	20
Scenario 3	5	2	10	20

Note. Team colors (Red, Blue) were randomized across participants. All three experiments used the same stimuli.

2), the two attacks occurred at the same time, meaning that the two attacks should be defended by different teams. Performance was again determined by the sum of the two selected competences. In the weakest-link structure (Experiment 3), the attacks occurred at the same time, but performance was determined by the weaker of the two assigned defense competences. This meant that performance was determined by how well this less competent team could hold its ground.

Each experiment consisted of three starting competence setups that defined the distribution of initial defense competences. In Setup 1, one team was better in both dimensions than the other. For example, in one version of the experiment,2 the Red team was better than the Blue in both defenses, and had a larger comparative advantage over the Blue team in air, but the Blue team was slightly better in air than land. In Setup 2, one team was better in one dimension, but equal in the other. For example, both teams were equally good in land, but Red team was better in air. In Setup 3, each team specialized in a different dimension: Red was better in land and Blue was better in air.

2.2. Theoretical framework

We first describe how we model the mind of individual teams (i.e., agents)—how they decide effort when they are assigned a training task—then describe how we model participants’ decisions as trainers. We denote the two military teams as agents X and Y, who each has a competence level in two dimensions (land defense and air defense) denoted as a and b.

2.2.1. Agents’ minds

In each round of training, an agent with competence c in the trained dimension is assigned a task of difficulty d. The agent chooses how much effort e ∈ [0, 1] to exert—which indicates what proportion of their competence is applied to the task—by maximizing expected utility

(Xiang et al., 2023; Xiang, 2026):

$$U(e) = rP(\text{success}|c, e, d) - \alpha e, \quad (1)$$

where r is the reward for succeeding in the training task and α is a scalar that indicates the cost of exerting effort (agents with larger α are lazier). According to this formalism, the agent may choose not to exert effort if the reward is small, or if the agent is lazy or unlikely to succeed. For simplicity, we assume that success follows a deterministic threshold rule that has been used in past work (Xiang, Landy, Cushman, Vélez, & Gershman, 2023, 2025):

$$P(\text{success}|c, e, d) = \mathbb{I}[ce \geq d], \quad (2)$$

where $\mathbb{I}[\cdot]$ is an indicator function that takes value 1 when the argument is true and value 0 otherwise. In other words, $P(\text{success}|c, e, d)$ is 1 when the applied competence (ce) exceeds the task difficulty d , and 0 if $ce < d$.

The optimal effort is given by:

$$e^* = \underset{e}{\operatorname{argmax}} U(e). \quad (3)$$

Following Xiang et al. (2024), we make the assumption that competence increases in proportion to effort:

$$c' = c + \rho e, \quad (4)$$

where $\rho \in [0, 1]$ is the learning rate. Thus, training choices reflect a balance between the team's current competence, the difficulty of the training task, and the anticipated effort they would exert: while a harder training task may require more effort and thus yield more improvement, a task that is too difficult would produce no effort, and thus no gains.³

2.2.2. Trainer decisions

In each of the three rounds of training, participants acted as a trainer who decided which team to train (X or Y), on which dimension (a or b), and at what difficulty. Each possible training trajectory i is a sequence of these training choices (who, what, and how hard) that starts with the pre-training competences $\mathbf{c} = \{c_{Xa}, c_{Xb}, c_{Ya}, c_{Yb}\}$ and ends with a set of post-training competences $\mathbf{c}'(i) = \{c'_{Xa}, c'_{Xb}, c'_{Ya}, c'_{Yb}\}$. The trainer assigns an action value $Q(i)$ to each training trajectory and selects among them probabilistically:

$$P(i) \propto \exp[\beta Q(i)], \quad (5)$$

where β is an inverse temperature parameter that controls the stochasticity of the choice. A trainer with β close to infinity would deterministically choose the training trajectory with the highest value, whereas a trainer with $\beta = 0$ would favor all trajectories equally. During the execution phase, the trainer assigns teams to the two attacks according to their execution values:

$$P(A) \propto \exp[\beta V(A)], \quad (6)$$

where A is a role assignment and $V(A)$ is the value of choosing assignment A , based on the post-training competence vector $\mathbf{c}'$. The models (summarized in the next section and detailed

in the Materials and Methods section) differ only in their choices of training and execution values.

Importantly, many of the simplifying assumptions here are not essential to the framework, which can be straightforwardly extended to accommodate different dynamics, such as allowing simultaneous training of both teams, discounting future training outcomes, or adopting more realistic learning rules (Xiang & Gershman, 2025).

2.2.3. Model formalization

We consider six models. Inspired by model-based approaches to sequential decision problems (Daw & Dayan, 2014; Huys et al., 2015), the *Planning model* plans training decisions to optimize final task performance in the execution phase. The model flexibly adjusts its training decisions according to the demands of the task structures (dimension-based, agent-based, or weakest-link), by anticipating which team would be assigned to defend each attack. The alternative models, on the other hand, apply the same heuristic rules across all task structures. First, the *Exploitation model* prioritizes training and deploying the already competent team. This is consistent with work showing how people prefer to invest in the currently higher-return options, thereby amplifying the advantages of initially stronger members (March, 1991; Merton, 1968). In contrast, people may also prioritize learning: human capital theory emphasizes a focus on future returns, whereby individuals invest in areas with the greatest expected improvements in productivity (Becker, 1964). Accordingly, the *Learning model* seeks to maximize total competence gain. Additionally, people also have concerns for fairness. Specifically, we distinguish between two types of fairness: equality, which is action-oriented, provides the same resources to everyone, whereas equity, which is outcome-oriented, provides resources in a way that reduces disparities between people (Minow, 2021). Thus, the *Equality model* expresses no preference among training options, giving all teams equal treatment and provides a baseline that tests whether the behavioral patterns deviate from random choice. By contrast, the *Equity model* prioritizes bringing the weakest closer to the others, effectively giving extra support to those who are behind. Lastly, prior work has also argued that generalist-or diverse-skilled profiles can be advantageous in uncertain environments since they support adaptability and cross-domain connections (Epstein, 2021). This thus motivates the *General model* that promotes balanced competence sets within each team.

During the training phase, the models chose a training strategy i associated with value $Q(i)$. Training resulted in post-training competences $\mathbf{c}'(i) = \{c'_{Xa}, c'_{Xb}, c'_{Ya}, c'_{Yb}\}$. In the execution phase, the models chose a role assignment A associated with value $V(A)$. In both phases, the choice rule was stochastic and governed by the softmax function (Eqs. 5 and 6) with inverse temperature β .

Planning model: The Planning model trains and deploys teams based on the task structure and goal.

Dimension based (Experiment 1): trains teams for the highest individual land and air defense competences.

Training value:

$$Q(i) = \max(c'_{Xa}, c'_{Ya}) + \max(c'_{Xb}, c'_{Yb}) \quad (7)$$

Execution value:

$$V(X \rightarrow a) = c'_{Xa}, \quad V(X \rightarrow b) = c'_{Xb}, \quad V(Y \rightarrow a) = c'_{Ya}, \quad V(Y \rightarrow b) = c'_{Yb} \quad (8)$$

Agent based (Experiment 2): trains teams for the highest individual land and air defense competences, while keeping in mind that each team will only handle one type of attack.

Training value:

$$Q(i) = \max(c'_{Xa} + c'_{Yb}, c'_{Xb} + c'_{Ya}) \quad (9)$$

Execution value:

$$V(X \rightarrow a, Y \rightarrow b) = c'_{Xa} + c'_{Yb}, \quad V(X \rightarrow b, Y \rightarrow a) = c'_{Xb} + c'_{Ya} \quad (10)$$

Weakest link (Experiment 3): trains teams to maximize the lower of the two defense competences, while keeping in mind that each team will only handle one type of attack.

Training value:

$$Q(i) = \max(\min(c'_{Xa}, c'_{Yb}), \min(c'_{Xb}, c'_{Ya})) \quad (11)$$

Execution value:

$$V(X \rightarrow a, Y \rightarrow b) = \min(c'_{Xa}, c'_{Yb}), \quad V(X \rightarrow b, Y \rightarrow a) = \min(c'_{Xb}, c'_{Ya}) \quad (12)$$

Exploitation model: The Exploitation model prioritizes training and deploying the more competent team.

Training value:

$$Q(i) = \max(c'_{Xa} + c'_{Xb}, c'_{Ya} + c'_{Yb}) \quad (13)$$

Execution value:

$$V(X \rightarrow a) = c'_{Xa}, \quad V(X \rightarrow b) = c'_{Xb}, \quad V(Y \rightarrow a) = c'_{Ya}, \quad V(Y \rightarrow b) = c'_{Yb} \quad (14)$$

Learning model: The Learning model trains teams to maximize the total competence gain. This principle also applies to deploying teams, but because it is unclear how much each team will improve during execution, the model approximates it by using the two training difficulties d_1 and d_2 . Thus, teams are assigned roles that have higher expected competence gain.

Training value:

$$Q(i) = c'_{Xa} + c'_{Xb} + c'_{Ya} + c'_{Yb} \quad (15)$$

Execution value:

$$\begin{aligned} V(X \rightarrow a) &= \max\left(\frac{c'_{Xa}}{d_1}, \frac{c'_{Xa}}{d_2}\right), & V(X \rightarrow b) &= \max\left(\frac{c'_{Xb}}{d_1}, \frac{c'_{Xb}}{d_2}\right) \\ V(Y \rightarrow a) &= \max\left(\frac{c'_{Ya}}{d_1}, \frac{c'_{Ya}}{d_2}\right), & V(Y \rightarrow b) &= \max\left(\frac{c'_{Yb}}{d_1}, \frac{c'_{Yb}}{d_2}\right) \end{aligned} \quad (16)$$

Equity model: The Equity model trains teams to minimize the competence difference between the two teams, and prioritizes deploying the less competent team.

Training value:

$$Q(i) = -|(c'_{Xa} + c'_{Xb}) - (c'_{Ya} + c'_{Yb})| \quad (17)$$

Execution value:

$$V(X \rightarrow a) = -c'_{Xa}, \quad V(X \rightarrow b) = -c'_{Xb}, \quad V(Y \rightarrow a) = -c'_{Ya}, \quad V(Y \rightarrow b) = -c'_{Yb} \quad (18)$$

Equality model: The Equality model is indifferent across all options.

Training value:

$$Q(i) = 0 \quad \text{for all } i \quad (19)$$

Execution value:

$$V(A) = 0 \quad \text{for all } A \quad (20)$$

General model: The General model trains and deploys teams to be more versatile and less specialized. This means that it prioritizes reducing the difference between the two defense competences of each team.

Training value:

$$Q(i) = -(|c'_{Xa} - c'_{Xb}| + |c'_{Ya} - c'_{Yb}|) \quad (21)$$

Execution value:

$$\begin{aligned} V(X \rightarrow a) &= c'_{Xb} - c'_{Xa}, & V(X \rightarrow b) &= c'_{Xa} - c'_{Xb}, & V(Y \rightarrow a) \\ &= c'_{Yb} - c'_{Ya}, & V(Y \rightarrow b) &= c'_{Ya} - c'_{Yb} \end{aligned} \quad (22)$$

2.3. Model fitting and comparison

For each computational model, we fit the inverse temperature parameter β that determines stochasticity in the softmax choice function in Eq. 5. We estimated β separately for each participant using maximum-likelihood estimation over that participant's training trajectories. For each model, we performed a grid search over $\beta \in [0, 10]$ in increments of 0.5, generating model-predicted probabilities for all possible training trajectories (i.e., the sequence of training decisions). The log-likelihood for each β was computed by summing the log probabilities assigned to the participant's observed trajectory, and the value of β that maximized this log-likelihood was retained as the participant-specific parameter estimate. For consistency and for testing the predictive power of the model fits, the same β values from fitting the training data were reused in the execution phase.

To assess how well each model captures participants' training and execution behavior, we compared individual participants' training and execution strategies to the model predictions using random-effects Bayesian model selection (Rigoux, Stephan, Friston, & Daunizeau, 2014). For each participant, we evaluated the fit of each model using $\text{BIC} = k \ln(n) - 2 \ln(L)$,

where k is the number of free parameters ($k = 1$ for training, $k = 0$ for execution because it used the same β values from training), 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 $-0.5 \times \text{BIC}$, 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 (Bishop, 2006). Across the three experiments, the Planning model achieved a protected exceedance probability of close to 1.00 for both training and execution, indicating that it is the most likely model in the population.

3. Results

3.1. People assigned roles to collaborators based on task demands and relative competences

Across all conditions, before engaging in training or assignment, participants first saw the starting competences of each team, which allowed them to form expectations about how each team would ultimately be deployed, given the task demands and the upcoming training opportunity. Their eventual execution choices thus reflected these earlier anticipations about who would be assigned to defend against which type of attack. To quantify the effects, we analyzed these patterns across experiments using separate Bayesian regression models. Each regression model predicted the probability of choosing each option (i.e., the proportion of trials in which each participant assigned a team to defend an attack), with the experiment index as the predictor to compare across task structures. Because the dependent variable was computed by aggregating across the three scenarios in each competence setup, each participant contributed only one observation. Therefore, the regression models did not include random effects.⁴ For all regression models, we used the default priors and report the posterior mean coefficient b and the 95% credible interval CrI, where an effect is credible if the CrI does not include 0.

As shown in Fig. 2, participants flexibly adapted their role assignments in the execution phase to the demands of each task structure. In the first competence setup (Fig. 2, left column), the Red⁵ team was more competent than the Blue team in both land and air, with a larger comparative advantage over the Blue team in air, but the Blue team was slightly better in air than land. Participants in Experiment 1 assigned the Red team to defend both land (mean $P = 0.95$) and air attacks (mean $P = 0.97$), aiming to maximize total defense. By contrast, when the two attacks needed to be handled by different teams in Experiment 2, participants assigned the Red team to defend air (mean $P = 0.78$) and assigned the Blue team to defend land (mean $P = 0.76$). This corresponded to a decrease in Red team being assigned land defense in Experiment 2 ($b = -0.71$, 95% CrI $[-0.76, -0.67]$), and an increase in Blue team being assigned land defense ($b = 0.71$, 95% CrI $[0.66, 0.76]$), compared to Experiment 1. When the goal shifted to minimizing loss from the weaker defense in Experiment 3, most participants maximized the weaker of the two assignments, which would fall on the Blue team, by assigning Blue team what it was better at—in this case, air defense (mean $P = 0.65$)—leaving the Red team to defend land attacks (mean $P = 0.65$). This led to a decrease in Blue team assigned to land ($b = -0.42$, 95% CrI $[-0.48, -0.35]$) and Red team assigned to

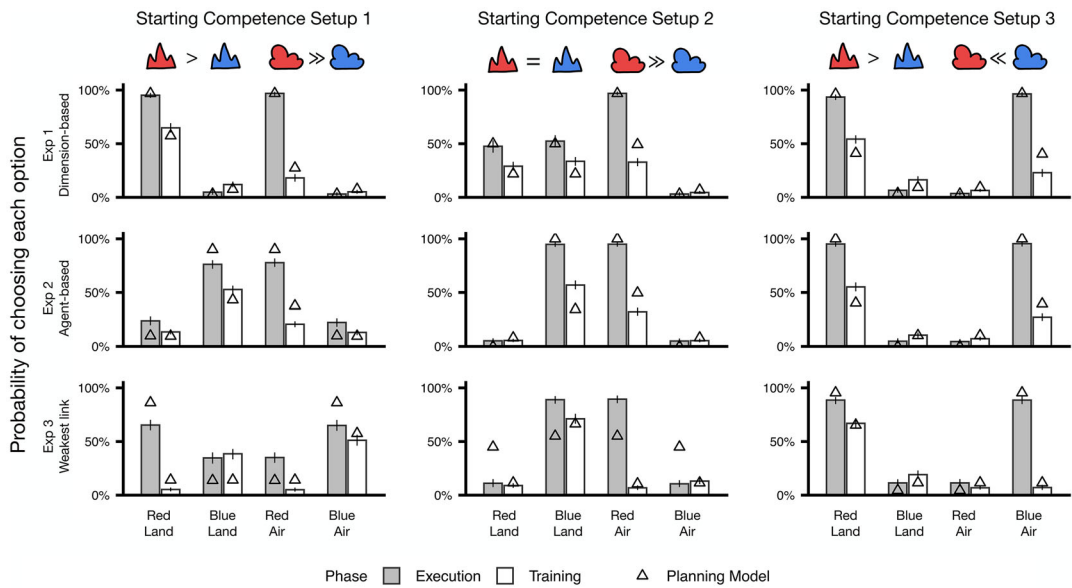


Fig. 2. Participants' training and execution decisions. Each panel corresponds to a starting competence setup (1–3). At the top of each column, the symbols correspond to the team-defense options on the x-axis (Red-Land, Blue-Land, Red-Air, Blue-Air): the mountain symbol denotes land defense competence, the cloud symbol denotes air defense competence, and the colors denote teams; these symbols illustrate the relative competence of the two teams. Each row corresponds to one of the three task structures (Experiments 1–3). Bars show the probability of choosing each option during the execution phase (gray) and training phase (white). Error bars indicate bootstrapped 95% confidence intervals. Overall, participants flexibly assigned agents to task dimensions to maximize performance under each task structure and trained teams for the roles they would later perform. These patterns were captured by the Planning model.

air ($b = -0.43$, 95% CrI $[-0.49, -0.36]$), and an increase in Red team assigned to land ($b = 0.42$, 95% CrI $[0.35, 0.48]$) and Blue team assigned to air ($b = 0.43$, 95% CrI $[0.36, 0.49]$).

The same principles applied to the other two competence setups as well. In the second competence setup (Fig. 2, middle column), the two teams were equally good in land but the Red team was better in air, participants assigned the Red team to defend air attacks in all experiments (mean $P = 0.97$ in Experiment 1, mean $P = 0.95$ in Experiment 2, mean $P = 0.89$ in Experiment 3). However, they were ambivalent about who should defend land attacks in Experiment 1 because the two teams were equally good in land (mean $P = 0.48$ for Red team and mean $P = 0.52$ for Blue team), but overwhelmingly assigned the Blue team to defend land attacks in Experiment 2 because the Blue team had to handle an attack (mean $P = 0.95$). These corresponded to pattern changes between Experiment 1 and Experiment 2, revealing a decrease in Red team assigned to land ($b = -0.42$, 95% CrI $[-0.48, -0.37]$) and an increase in Blue team assigned land ($b = 0.42$, 95% CrI $[0.37, 0.48]$) from Experiment 1 to Experiment 2. Experiment 3 concerned maximizing the lower of the two assigned competences. Because Blue-air began with the same competence level as Red-land and Blue-land, it was marginally better to train Blue-land and let it serve as the lower of the two assigned

competences. Splitting training across Red-land and Blue-air would still leave one of them as the limiting competence. In line with this, compared to Experiment 1, participants in Experiment 3 were more likely to assign the Blue team to defend land attacks (mean $P = 0.89$), revealing a decrease in Red team assigned to land ($b = -0.18$, 95% CrI $[-0.22, -0.15]$), and an increase in Blue team assigned land ($b = 0.18$, 95% CrI $[0.15, 0.21]$).

In the third competence setup (Fig. 2, right column), when the Red team was better in land and the Blue team was better in air, participants consistently assigned the Red team to defend land (mean $P = 0.93$ in Experiment 1, mean $P = 0.95$ in Experiment 2, mean $P = 0.89$ in Experiment 3) and Blue team to defend air (mean $P = 0.96$ in Experiment 1, mean $P = 0.95$ in Experiment 2, mean $P = 0.89$ in Experiment 3).

3.2. People made training decisions according to anticipated roles

Before the execution phase, participants completed three rounds of training to improve the teams' defense competences. In each round, participants chose to train one *team* (Red or Blue) on one type of *defense competence* (land or air) using one of the two available *training tasks* with various difficulties. As shown in the training phase decisions in Fig. 2, across experiments, participants tended to train teams on competences that would be relevant for the execution phase, showing that their training decisions reflected their anticipation of the roles each team would later play in execution. To maximize performance, training was allocated more to teams that could improve faster with the available training difficulties, with the exception of Experiment 3, where training was most helpful on the weaker team. As with execution data, here we analyzed the training patterns across experiments using separate Bayesian regression models. Each regression model predicted the proportion of trials in which participants trained a team on a defense dimension (e.g., Red team-land), with the experiment index as the predictor to compare across task structures.

In the first starting competence setup (Fig. 2, left column), since the Red team was expected to handle land in Experiment 1 but not in Experiment 2, participants trained Red on land much less in Experiment 2 than in Experiment 1 ($b = -0.51$, 95% CrI $[-0.57, -0.46]$). And since the Blue team was expected to handle air in Experiment 3 but not in Experiment 2, participants trained Blue on air much more in Experiment 3 than in Experiment 2 ($b = 0.38$, 95% CrI $[0.32, 0.44]$). In the second competence setup (Fig. 2, middle column), because the Blue team was more preferred for land attacks in Experiment 2 compared to Experiment 1, participants trained Blue on land more in Experiment 2 than in Experiment 1 ($b = 0.23$, 95% CrI $[0.17, 0.29]$). Furthermore, even though Experiments 2 and 3 both assigned Red to air and Blue to land, because Experiment 3's goal emphasized maximizing the weaker competence of the two and Blue was weaker, Blue was trained on land more in Experiment 3 than in Experiment 2 ($b = 0.14$, 95% CrI $[0.08, 0.20]$). Finally, in the third competence setup (Fig. 2, right column), even though the expected role assignment was the same across all experiments (Red team to land and Blue team to air), because Experiment 3 maximized the weaker defense competence and Red's land competence was worse than Blue's air compe-

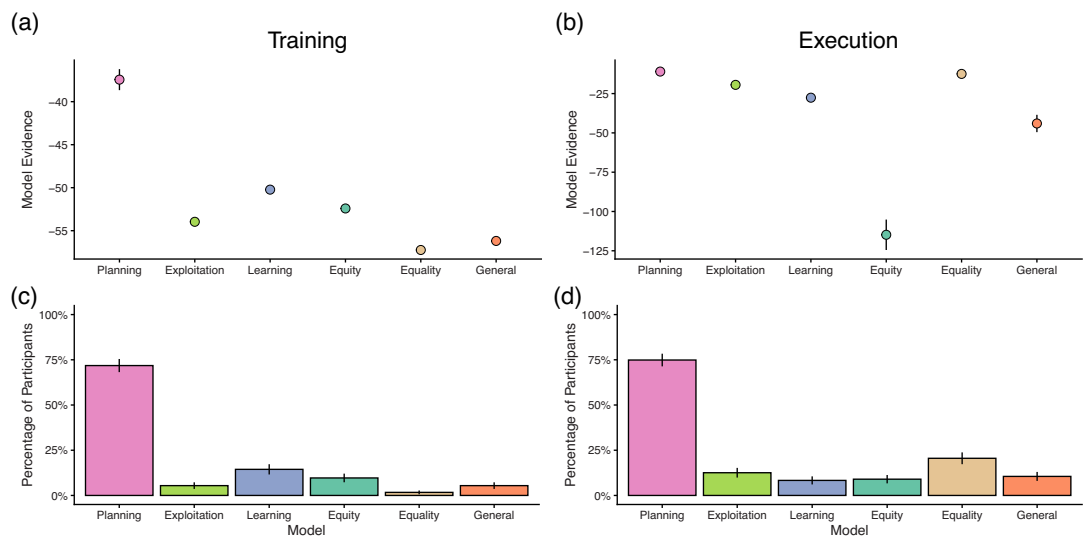


Fig. 3. (A, B) Approximate model evidence (log marginal likelihood) for training and execution, aggregated across experiments. (C, D) Proportion of participants whose behavior was best fit by each model during training and execution, aggregated across experiments. Error bars indicate bootstrapped 95% confidence intervals. Across both phases, the Planning model has the highest total model evidence and best explains training decisions for the majority of participants.

tence, participants trained Red on land more in Experiment 3 than in Experiment 1 ($b = 0.06$, 95% CrI [0.03, 0.09]) or in Experiment 2 ($b = 0.12$, 95% CrI [0.05, 0.18]).

In short, across starting competence setups, participants directed training to where it would be the most beneficial to the success of the collaboration, supporting the anticipated roles in execution.

3.3. The Planning model captured participants' decisions better than heuristic models

We compared participants' training and execution decisions to the six computational models defined in the Theoretical framework section. As shown in Fig. 2, the Planning model aligned well with participants' data. It is also the only model that captured the different training patterns observed across experiments; all heuristic models predicted the same behavior across experiments (see Figs. S1–S3).

Quantitatively, we conducted a random-effects Bayesian model selection (Rigoux et al., 2014), a standard approach for comparing models that evaluates how well each model fits individual participants' data. The Planning model had a protected exceedance probability of close to 1.00, indicating that it is decisively the most likely model in the population. Fig. 3A, B shows the distribution of model evidence across participants; the Planning model had the highest total model evidence for both training and execution. In line with this, Fig. 3C,D shows the proportion of participants whose behavior was best explained by each model; the

Planning model best accounted for the largest proportion of participants. Furthermore, consistent with these results, the distribution of fitted inverse temperature parameters β shows higher average β value under the Planning model. In contrast, β values were largely near zero for the heuristic models (except the Learning model). With a β of 0, all choices essentially have the same value; thus, this pattern suggests that the model predictions fail to meaningfully explain participants' behavior (see Fig. S4).⁶

4. Discussion

Division of labor is central to collaboration, but prior research has largely treated collaborators' competences as fixed inputs to be matched to task demands. Here, we show that people do not simply allocate tasks based on current competences—they actively develop those competences in ways that support the roles they expect them to play. These behavioral patterns were best explained by a forward-looking Planning model, which considers how training will impact future performance, and outperformed a range of simpler heuristics.

The Planning model formalizes prospective reasoning by simulating all possible training trajectories and selecting those that best satisfy anticipated task demands. This formulation is intended as a way to formalize the nonmyopic nature of people's decisions, and people almost certainly do not enumerate the full combinatorial space of trajectories. Instead, they may approximate planning through structured simplifications—for example, by first narrowing the space of possibilities based on expectations about the likely division of labor, and then evaluating only the training paths relevant to those anticipated roles. Importantly, this framework accounts for cases where training *alters* the division of labor itself. If we only considered starting competences, in competence setup 3 of Experiment 3, the two divisions of labor—Red-air and Blue-land versus Red-land and Blue-air—would be equally preferred, but participants clearly preferred one over the other (see Fig. 2), suggesting that training considerations inform the choice of how labor is divided. Consider another simple example where person A begins slightly more competent than person B but person B improves more quickly. With a long enough horizon, a forward-looking planner would invest in B and assign them the relevant role, whereas accounts that divide labor solely based on initial competences would assign the role to person A. This ability to capture how training reshapes future roles underscores why a prospective framework is essential for understanding division of labor.

Our findings connect to broader themes in predictive cognition, which emphasize that the mind is oriented toward anticipating future states to guide behavior, including in social contexts where people organize knowledge about others to predict their future actions (Clark, 2013; Tamir & Thornton, 2018). While prior work has focused on predicting sensory inputs or others' mental states and stable traits, our results suggest that people also form prospective inferences about how others' competences will evolve and use these inferences to guide their actions. This perspective also complements research on team cognition, which emphasizes that productive teams rely on structural supports external to individual minds, such as material artifacts and institutional roles (Fiore et al., 2010; Hutchins, 1991, 1995). Rather

than assuming static competences and roles, our work highlights the importance of anticipating how they will change. Incorporating this anticipation into existing accounts of predictive cognition and team cognition will help expand their explanatory power.

Our results show that division of labor is not inherently preferred in collaboration; its benefits depend on the underlying structure of competences and task demands. In Experiment 1, specialization clearly paid off only when each team was comparatively stronger in a different competence dimension. When the teams were equal in one dimension or when one team dominated both, dividing labor provided little advantage over assigning both roles to one team. This pattern aligns with prior work demonstrating that the value of dividing labor depends critically on the structure and size of the team (Cooper et al., 2021; Powell, 2000). The current study adds to this picture by showing that people flexibly adopt or forgo division of labor depending on whether it supports the goals of the collaboration. An open question is how these findings extend to settings with tighter, moment-to-moment interdependence among collaborators—for example, the collaboration between an architect and a builder (McCarthy, Hawkins, Wang, Holdaway, & Fan, 2021)—where collaborators with complementary strengths must act jointly on the *same* task.

A limitation of the current experiments is that the task environment is simplified: training outcomes were deterministic and fully observable, and future task demands were specified in advance. However, collaboration in naturalistic environments is often less transparent: people must infer others' competences and growth from noisy signals or incomplete information (Mercier, Morin, Mercier, & Quillien, 2026; Xiang, Gershman, & Gerstenberg, 2026), and collaboration often involves unclear and changing goals and time horizons, as well as overlapping roles. Under these conditions, future scenarios are not well-defined, and thus the feasibility of planning around specific goals becomes limited. As a result, people may rely on simpler heuristics that are not tailored to a specific future scenario, but more generally guide progress in a direction that benefits the collaboration (e.g., a heuristic that prioritizes maximal learning). These heuristics may be particularly preferred when the decision problem becomes more complex or involves time pressure (Payne, Bettman, & Johnson, 1993). An open question is under what contexts do simpler heuristics better explain human data in this kind of planning problem and how people determine which heuristics are most relevant. More fundamentally, future research can examine how collaborative goals emerge in the first place, and how people form an anticipation of future division of labor in the absence of such clear goals: do people train all collaborators equally (much like the Equality model) until a specific goal is in place, or do people choose training trajectories that would be most flexible to meet a range of goals? Beyond its specific formalism, the Planning model offers an explanation for why, for example, the Exploitation heuristic may be preferred when competence gains are negligible, or why the Equity heuristic may be favored under high team turnover.

Together, these findings indicate that people do more than match collaborators to roles—they develop collaborators in anticipation of those roles, and the development trajectories guide the role expectations. This perspective provides a starting point for understanding how

the plasticity of competence supports effective division of labor, and how humans cultivate and ultimately benefit from each other's evolving competences.

Acknowledgments

We thank Fiery Cushman and Shuze Liu for helpful discussions. This research was supported by the Kempner Institute for Natural and Artificial Intelligence, and a Polymath Award from Schmidt Sciences to S.J.G.

Notes

- 1 The other two advantages are saving the time lost in switching between tasks and facilitating the invention of machinery (Smith, 1937; Yang & Ng, 1998).
- 2 The colors of the teams were counterbalanced in the experiments.
- 3 For all experiments, we used a fixed learning rate $\rho = 1$, which meant that the teams' competence increased by the amount of effort they exerted. This information was available to participants. We also set $r = 10$ and $\alpha = 5$ for all experiments. These values were chosen so that the teams were motivated to exert effort, but the effort costs were not negligible. To avoid information overload, r and α were not shown to participants, but participants knew from the examples that the teams were sufficiently motivated, such that they would complete any task that is within their competence with the minimal effort needed, and give up on tasks that are beyond their competence. The specific values of r and α did not matter for our experiments; our examples clarified that $r > \alpha$, and any values that satisfy $r > \alpha$ would lead to the same result.
- 4 This was a minor deviation from the preregistration. We explain the deviation in detail in the Supplementary Material.
- 5 We randomized the color of the teams in the experiments.
- 6 Note that the Equality model by definition assigns all options the same value. Thus, the fitted beta values here do not provide information on the alignment between data and model.

References

- Becker, G. S. (1964). *Human capital: A theoretical and empirical analysis, with special reference to education*. New York: National Bureau of Economic Research.
- Bishop, C. M. (2006). *Pattern recognition and machine learning*. Springer.
- Clark, A. (2013). Whatever next? Predictive brains, situated agents, and the future of cognitive science. *Behavioral and Brain Sciences*, 36(3), 181–204.
- Cooper, G. A., Frost, H., Liu, M., & West, S. A. (2021). The evolution of division of labour in structured and unstructured groups. *Elife*, 10, e71968.
- Curioni, A. (2022). What makes us act together? On the cognitive models supporting humans' decisions for joint action. *Frontiers in Integrative Neuroscience*, 16, 900527.
- Daw, N. D., & Dayan, P. (2014). The algorithmic anatomy of model-based evaluation. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 369(1655), 20130478.

- Epstein, D. (2021). *Range: Why generalists triumph in a specialized world*. Penguin.
- Fiore, S. M., Rosen, M. A., Smith-Jentsch, K. A., Salas, E., Letsky, M., & Warner, N. (2010). Toward an understanding of macrocognition in teams: Predicting processes in complex collaborative contexts. *Human Factors*, 52(2), 203–224.
- Goldstone, R. L., Andrade-Lotero, E. J., Hawkins, R. D., & Roberts, M. E. (2024). The emergence of specialized roles within groups. *Topics in Cognitive Science*, 16(2), 257–281.
- Haeussler, C., & Sauermann, H. (2016). The division of labor in teams: A conceptual framework and application to collaborations in science. Technical report, National Bureau of Economic Research.
- Haeussler, C., & Sauermann, H. (2020). Division of labor in collaborative knowledge production: The role of team size and interdisciplinarity. *Research Policy*, 49(6), 103987.
- Hutchins, E. (1991). The social organization of distributed cognition. In L. B. Resnick, J. M. Levine, & S. D. Teasley (Eds.), *Perspectives on socially shared cognition* (pp. 283–307). American Psychological Association.
- Hutchins, E. (1995). *Cognition in the wild*. MIT Press.
- Huys, Q. J., Eshel, N., O’Nions, E., Sheridan, L., Dayan, P., & Roiser, J. P. (2012). Bonsai trees in your head: How the pavlovian system sculpts goal-directed choices by pruning decision trees. *PLoS Computational Biology*, 8, e1002410.
- Huys, Q. J. M., 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.
- Liu, S., Xiang, Y., & Gershman, S. J. (forthcoming). Probabilistic forecasting guides dynamic decisions. *Psychological Review*.
- March, J. G. (1991). Exploration and exploitation in organizational learning. *Organization Science*, 2(1), 71–87.
- McCarthy, W. P., Hawkins, R. D., Wang, H., Holdaway, C., & Fan, J. E. (2021). Learning to communicate about shared procedural abstractions. *arXiv preprint arXiv:2107.00077*.
- Meder, B., Nelson, J. D., Jones, M., & Ruggeri, A. (2019). Stepwise versus globally optimal search in children and adults. *Cognition*, 191, 103965.
- Mercier, M., Morin, O., Mercier, H., & Quillien, T. (2026). Who knows what? Bayesian competence inference guides knowledge attribution and information search. *Cognition*, 273, 106533.
- Merton, R. K. (1968). The Matthew effect in science: The reward and communication systems of science are considered. *Science*, 159(3810), 56–63.
- Minow, M. (2021). Equality vs. equity. *American Journal of Law and Equality*, 1, 167–193.
- Payne, J. W., Bettman, J. R., & Johnson, E. J. (1993). *The adaptive decision maker*. Cambridge University Press.
- Powell, S. G. (2000). Specialization, teamwork, and production efficiency. *International Journal of Production Economics*, 67(3), 205–218.
- Rieskamp, J., & Hoffrage, U. (1999). When do people use simple heuristics, and how can we tell. In G. Gigerenzer, P. M. Todd, & the ABC Research Group (Eds.), *Simple Heuristics that Make Us Smart* (pp. 141–167). Oxford 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* (pp. 3–16). New York: Modern Library.
- Tamir, D. I., & Thornton, M. A. (2018). Modeling the predictive social mind. *Trends in Cognitive Sciences*, 22(3), 201–212.
- 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.
- Wahn, B., & Schmitz, L. (2023). Labor division in collaborative visual search: A review. *Psychological Research*, 87(5), 1323–1333.
- Xiang, Y. (2026). *Cognitive foundations of collaboration*. PhD thesis, Harvard University.

- Xiang, Y., & Gershman, S. (2025). Modeling intrinsic motivation as reflective planning. In *Proceedings of the Annual Meeting of the Cognitive Science Society*, volume 47.
- Xiang, Y., Gershman, S. J., & Gerstenberg, T. (2026). A signaling theory of self-handicapping. *Cognition*, 266, 106288.
- Xiang, Y., Landy, J., Cushman, F. A., Vélez, N., & Gershman, S. J. (2023). Actual and counterfactual effort contribute to responsibility attributions in collaborative tasks. *Cognition*, 241, 105609.
- Xiang, Y., Landy, J., Cushman, F. A., Vélez, N., & Gershman, S. J. (2025). People reward others based on their willingness to exert effort. *Journal of Experimental Social Psychology*, 116, 104699.
- Xiang, Y., Vélez, N., & Gershman, S. J. (2023). Collaborative decision making is grounded in representations of other people's competence and effort. *Journal of Experimental Psychology: General*, 152(6), 1565.
- Xiang, Y., Vélez, N., & Gershman, S. J. (2024). Optimizing competence in the service of collaboration. *Cognitive Psychology*, 150, 101653.
- Yang, X., & Ng, S. (1998). Specialization and division of labour: A survey. In K. J. Arrow, Y.-K. Ng, & X. Yang (Eds.), *Increasing Returns and Economic Analysis* (pp. 3–63). Springer.

Supporting Information

Additional supporting information may be found online in the Supporting Information section at the end of the article.

Data S1