

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Discovering and transmitting abstract knowledge over generations

Permalink

<https://escholarship.org/uc/item/753715pj>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Prystawski, Ben

Frank, Michael C.

Goodman, Noah

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Discovering and transmitting abstract knowledge over generations

Ben Prystawski (benpry@stanford.edu)¹, Michael C. Frank¹ & Noah D. Goodman^{1,2}

¹Department of Psychology, Stanford University

²Department of Computer Science, Stanford University

Abstract

The complexity of human culture depends on people’s ability to discover and transmit abstract knowledge. Studying this ability is crucial to understanding humans’ distinctive place among species, but current experimental paradigms focus on the cultural transmission of specific, concrete facts rather than generalizable abstract knowledge. In this paper, we develop a crafting game paradigm to study how people discover abstract knowledge and transmit it via language. We compared individuals playing this game for 40 rounds to chains of four participants playing for 10 rounds each and passing messages to each other sequentially. The individuals performed significantly better over rounds, but the chains did not. Through simulations with language model agents and a follow-up experiment, we find substantial variation in the helpfulness of participants’ messages, which may explain the lack of consistent improvement in chains. The ability to learn selectively from the good messages may be essential for improvement over generations.

Keywords: Iterated learning; cultural evolution; language; crafting

Introduction

Humans owe our distinctive role in the world to the capacity to learn from each other and build up vast bodies of knowledge over generations (Boyd et al., 2011; Mesoudi, 2016; Tomasello, 1999). Human culture contains many concrete facts, like knowledge of which mushrooms are safe to eat. But culture also gives us abstract, generalizable principles. Learning a skill like sailing requires one to not only memorize specific sequences of actions, but to learn principles of navigation and how to respond to different weather conditions. Understanding how humans transmit abstract knowledge is essential to understanding the unparalleled complexity of human culture.

Cognitive science has sought to understand how the human mind enables complex culture by running experiments in which people learn from each other. Some studies analyze pairs of participants with one teacher and one learner (Chopra et al., 2019; Sumers et al., 2023). Other studies use *iterated learning*, where participants successively interact with a task, learn something about it, then transmit that knowledge to the next participant (Beppu & Griffiths, 2009; Smith et al., 2003). In some of these studies, participants write messages to each other in language, while in others they learn by observing others’ actions (Martinez et al., 2024; Witt et al., 2024).

Prior studies on teaching and iterated learning have mostly

used simple tasks like multi-armed bandits, category learning, and function learning, where the results can connect directly to an existing body of theory. However, the simplicity of these tasks prevents them from being a good setting to study the transmission of abstract knowledge. A few studies have used tasks with abstract structure (e.g. Ham et al., 2025; Tessler et al., 2021). However, these tasks either use simple abstractions or use video game environments that are so complex it is difficult to analyze human behavior in a fine-grained way.

Developing a task that combines rich abstract structure with tractability of analysis is valuable to study the transmission of abstractions experimentally. For instance, mathematical and computational models make predictions about how cultural learning is shaped by transmission fidelity (Lewis & Laland, 2012; Prystawski et al., 2023, 2025), biases in individual learning (Thompson & Griffiths, 2021), and the ability to integrate linguistic guidance with direct experience (Colas et al., 2025). Many of these models are based in simple tasks, and it is an open question whether their predictions hold for more complex tasks: after all, a model that explains how people learn what mushrooms are safe to eat might not explain the development of mathematics.

Crafting games, where players combine items from an inventory to create new items, offer a combination of measurability and open-ended complexity. In these games, the space of actions available at any one time is limited (combining pairs of items), but players can create potentially infinitely many items over time. People find crafting games fun and are intrinsically motivated to do well in them, as evidenced by the popularity of games like *Little Alchemy* and *Infinite Craft*. *Little Alchemy* has been used to study how people explore (Brändle et al., 2022, 2023), and a crafting task has been used to study the effects of network structure and semantic knowledge on cultural evolution (Derex & Boyd, 2015; Yaman et al., 2025). Crafting involves recombining existing discoveries to produce new things, which is an essential part of human innovation in open-ended domains (Morgan & Feldman, 2025; Zhao et al., 2024, 2025).

In this paper, we develop a crafting game where items combine according to abstract rules and use it to study how people discover and transmit knowledge. We ran an iterated learning experiment in which chains of four participants played the game and passed messages to the next member of their chain, alongside “immortal individuals” (Tessler et al., 2021) who played the game for much longer. While participants’ scores

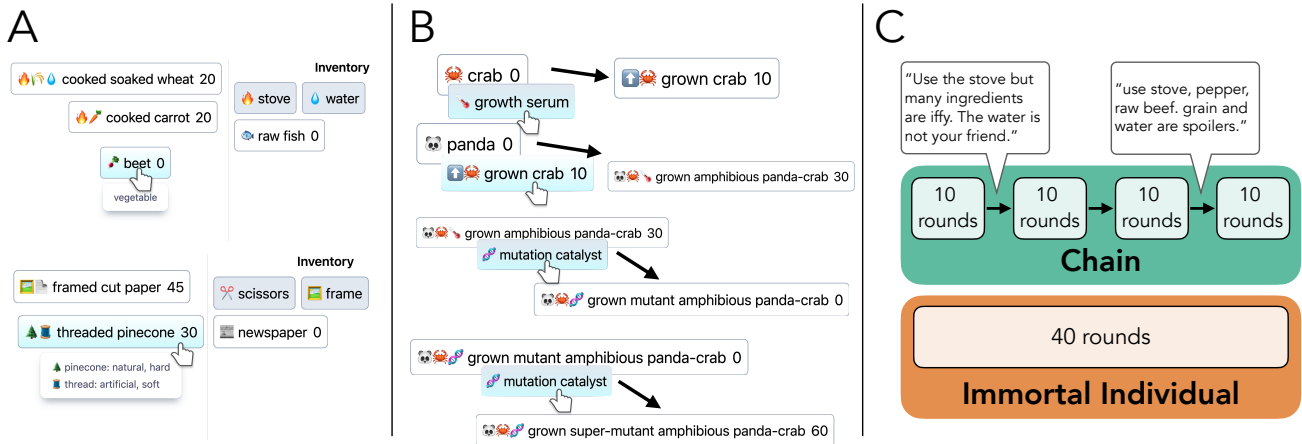


Figure 1: Overview of the crafting task and experimental paradigm. **A**: Two screenshots of the game environment. The participant has an inventory of tools (grey background) and ingredients (white background) which they can click and drag into the crafting area and combine. **B**: Example sequence of combinations in the animals domain. **C**: Illustration of the design for Experiment 1. Participants in the chain condition took part in a chain of 4 participants and played 10 rounds in each domain. They received a message from the previous chain member and wrote a message to the next. Participants in the immortal individual condition played in one domain for 40 rounds.

improved significantly over rounds in the immortal individual condition, they did not in the chain condition. However, some participants performed well at the task and wrote helpful messages. We develop a simulation framework using language model agents to evaluate the quality of messages. In a follow-up experiment, we validate the simulations of message value by comparing human performance given the messages identified as best and worst by our simulations. We find that our simulation framework generally predicts how valuable messages are to participants, but overestimates how well people do when given the best messages.

Experimental Paradigm

We developed an experimental paradigm in the context of a crafting game. Figure 1A shows an example of this game’s interface. In the game, participants have an inventory with two kinds of items: tools and ingredients. Each ingredient has a set of features, and ingredients have values between 0 and 100. The rules for how valuable an item is and how items combine are determined by the features. The participant starts each round of the game with two tools and four ingredients. Tools can be reused as much as the participant wants, but ingredients can be used only once. The participant can press a “Submit” button at any time after making at least one combination to end the round. When the round ends, the score is the value of the highest-value item in their inventory. Each round’s starting inventory has the same tools, but a different random set of starting ingredients. Randomized starting inventories mean that participants need to generalize from specific item combinations to learn the feature-based rules determining how items combine and how valuable they are.

The game has a drag and drop interface where the partici-

pant clicks and drags items from their inventory into a crafting area in the center of the screen. Then the participant can combine two items by dropping one on top of the other. An item’s features can be seen by hovering over the item with the mouse. The starting inventories are generated with constraints that ensure that the maximum achievable score in every round is exactly 100.

Each item has a name and an emoji. The starting items are hand-designed, but all other items’ names and emoji are assigned using a language model. This lets us create a vast space of possible items by writing functions that define how the features behave and instructions for the language model on how to name items. We used gpt-oss-20b for this task as we found it balanced low latency with accurate rule following (Agarwal et al., 2025). Figure 1B shows an example sequence of combinations and their results.

The general framework of this game lets us define different domains, each with different cover stories and different sets of starting ingredients, features, tools, and underlying rules. We designed four domains for our experiments: *cooking*, where the participant combines items in a kitchen to create delicious dishes; *decorations*, where the participant combines household items to make nice decorations; *animals*, where the participant combines and modifies animals to create hybrid creatures; and *potions*, where the participant combines ingredients to make potions.

Experiment 1

Experiment 1 was an iterated learning study designed with two goals. First, we wanted to validate the crafting game paradigm by comparing chains of participants learning and sending messages to each other against a baseline of “immortal

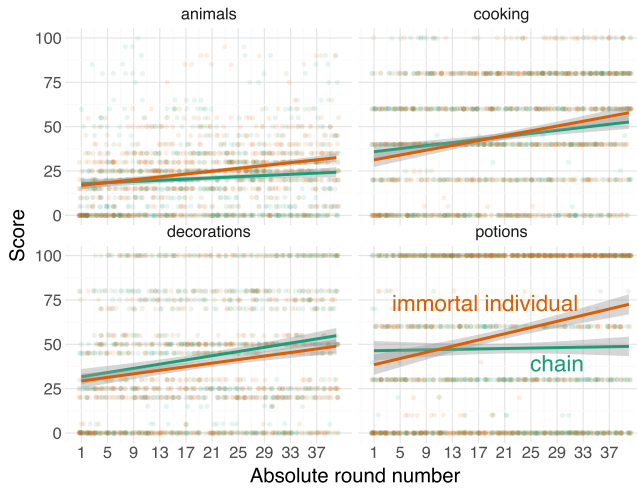


Figure 2: Performance by absolute round number for participants in the immortal individual and chain conditions of Experiment 1. Error ribbons show 95% confidence intervals.

individuals” who play for the same number of rounds as an entire chain. Second, we wanted to understand what people write in messages to each other and how helpful different kinds of messages are to their recipients. This experiment was implemented using the PsyNet framework (Harrison et al., 2020). Data and analysis code can be found at <https://github.com/benpry/transmitting-abstract-knowledge>.

Methods

Participants 169 participants completed the experiment on Prolific.¹ An additional 12 completed the post-experiment survey, but had timed out earlier in the experiment and had all of their data excluded. Participants were paid a base payment of \$10.14 and a bonus of up to \$1 (average 41¢) based on their performance. They took a median of 46 minutes to complete the study. In the chain condition, bonuses were half based on their individual performance and half based on the performance of the participants who received their messages.

Procedure Participants first read instructions explaining how the game interface worked and how they would be scored. Next, they played a simple practice game where the items were playing cards and there was only one tool. This practice game repeated with the same starting inventory until the participant earned a score of 80 or more.

After completing the practice round, participants were assigned to either the *immortal individual* condition or the *chain* condition. Participants in the chain condition played in four blocks, one for each domain, for 10 rounds each. In each block, the participant was randomly assigned to one of 20 chains for that block’s domain. Each chain consisted of four participants, and we refer to a participant’s position within

¹The preregistration for this experiment can be found at <https://osf.io/yr7jv/>.

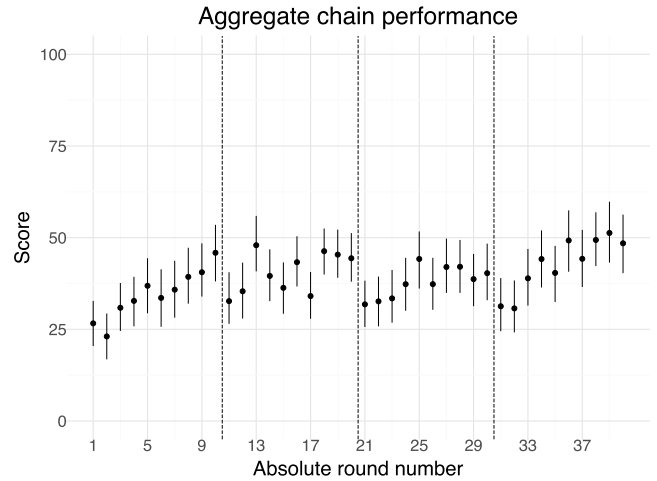


Figure 3: Performance of participants in the chain condition by absolute round number, aggregated across domains. Dashed lines denote transitions between generations. Error bars show bootstrapped 95% confidence intervals.

their chain as their “generation.” Participants read a message from the previous participant in their chain before playing the game and wrote a message to the next participant after playing. Participants in the immortal individual condition played in one domain for 40 rounds. These participants did not receive a message, but wrote a message after playing the game. Figure 1C illustrates the two conditions.

If a participant in the chain condition took longer than 20 minutes to complete any block, they timed out. That block, along with their previously-completed blocks, was marked as failed. If they returned to the experiment and completed more blocks, those blocks were not failed. Due to these constraints, we lost 24 blocks from participants who otherwise completed the experiment. Nine of our participants were recruited to make up for these lost blocks. We also excluded two blocks from participants in generation four who were accidentally assigned to the same domain twice.

Results

Overall learning We preregistered two main hypotheses. The first is that there would be a significant improvement in performance with absolute round number (number of rounds elapsed from the beginning of a chain, across all participants) in both conditions. The second is that participants would learn more slowly in the chain condition. We tested these hypotheses in a Bayesian mixed-effects model predicting score as a function of absolute round number and condition with random intercepts and slopes for each participant or chain.² We tested the first hypothesis using the marginal effect of

²The full model formula is $\text{score} \sim 1 + \text{absolute_round_num} + \text{condition}:\text{absolute_round_num} + (1 + \text{absolute_round_num} \mid \text{domain} / \text{participant_or_chain_id})$

absolute round number in each condition and the second using the interaction between condition and absolute round number. We report coefficient estimates with 95% credible intervals.

Figure 2 shows participants’ scores over rounds by domain and condition. We found a significant marginal effect of absolute round number in the immortal individual condition ($\beta = 0.60$ [0.03, 1.16]), but no significant effect in the chain condition ($\beta = 0.33$ [−0.23, 0.90]). The interaction between condition and absolute round number was significant ($\beta = 0.27$ [0.06, 0.48]). Participants exhibited significant learning in the immortal individual condition, but not the chain condition. These results align with our second hypothesis, but partially contradict our first.

We also ran a follow-up analysis on data in the chain condition in which we modeled learning over generations separately from learning within generations.³ This model also showed no significant improvement in scores over chains ($\beta = 2.29$ [−3.69, 8.68]).

There was dramatic variability in score by domain and by participant. The standard deviation on the random slopes was 0.19 by domain and 0.60 by participant or chain within domains, indicating a wide range of learning rates. While a chain might make steady progress over a few generations, a participant late in the chain who is not very engaged with the task – or who tries to learn the rules completely at the expense of achieving high scores – can lower the overall slope of learning progress. The potions domain also shows a particularly stark difference between chain and immortal individual conditions.

Figure 3 shows the average score for each round in the chain condition across all domains. Learning shows a sawtooth-shaped trajectory where each participant improves over rounds, but performance drops at the transitions between generations. While participants made progress, their progress was mostly lost in transmission to the next participant.

Message content Messages in the chain condition had an average of 26.64 words (SD: 20.99). To understand the content of the messages people wrote, we coded them along seven different psychological dimensions using a coding scheme inspired by the scheme used by Tessler et al. (2021). These dimensions are abstraction, concrete recipes, positive valence, negative valence, policy, dynamics, and avoidance. Abstraction refers to whether a message has advice that could apply to more than a specific set of items. Concrete recipes measures whether a message mentions specific items by name and describes what results from combining them. Positive valence codes for whether a message mentions things like winning, achieving high scores, and what the best strategy is. Negative valence codes for whether a message mentions losing, not being able to do better than a certain score, or being confused. Policy refers to whether a message has instructions for what the recipient should do. Dynamics codes for whether

³The formula for this model is $\text{score} \sim 1 + \text{chain_pos} + \text{round_num} + (1 + \text{chain_pos} \mid \text{domain} / \text{participant_or_chain_id})$

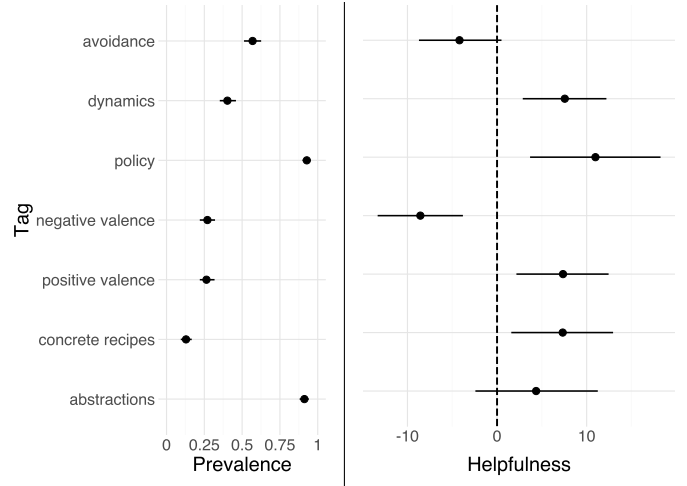


Figure 4: Prevalence and helpfulness of each tag for messages in Experiment 1. Error bars show bootstrapped 95% confidence intervals.

the message has descriptions of how the game works, including descriptions of what produces high and low-value items and how tools influence features. Finally, avoidance describes whether the message says what not to do.

The first author manually coded a subset of 105 messages, then used an automated language model-based coding pipeline to code the remaining messages. 20 of the human-coded messages were selected as in-context examples for the language model. 45 of the messages were used to evaluate agreement with human judgments for development purposes, including tweaking the instructions to the model and selecting the best model. The remaining 40 were held out for a final evaluation. We chose DeepSeek V3.2 for coding, as it agreed most with human ratings on the development set (Liu et al., 2025). The model’s tags agreed with human tags an average of 86% of the time on the evaluation set, with the lowest-agreement tag (dynamics) having an agreement of 70%.

Figure 4 (left) shows the proportion of messages in the chain condition that contain each tag. Most messages contain at least some abstraction and mention policy. Only about 25% of messages have positive or negative valence, and very few mention concrete recipes. A minority of messages talk about dynamics, and about half mention avoidance.

Simulations of message value

How useful are messages for the people who read them? If a participant does well after reading a message, it could be because the participant was a good learner or because the message was helpful – the two are completely confounded in our paradigm. To address this issue, we developed a simulation framework that uses language model agents to play the game given human messages. The agent is given instructions similar to the instructions we provided human participants. On each timestep of the game, it sees a text representation of the current game state which includes the inventory and the

features and values of all the items. The model responds with a JSON object containing two fields: “reasoning” and “action”. “Reasoning” contains a one- or two-sentence rationale for choosing an action and “action” contains either the names of two items to combine or the word “submit”. The agent can play for several rounds, with previous rounds kept in context. Figure 5 shows an example round in this framework.

For each human message from Experiment 1, we ran an agent that saw the message and measured its performance over five rounds of gameplay. We ran this agent 10 independent times for each message in order to produce a reliable estimate. We define the *message value* as the average performance of the agent over all 50 rounds of this simulation. All of our simulations used Gemini 2.5 Flash as the language model (Comanici et al., 2025).

We first compared message value to the performance of senders and recipients of each message in Experiment 1. We found significant Pearson correlations between message value and the performance of both the senders ($r = 0.57, p < 0.001$) and recipients ($r = 0.50, p < 0.001$) in the chain condition. These correlations are both stronger than the correlation between the performances of senders and their recipients ($r = 0.33, p < 0.001$).

We then analyzed the content of the messages along the dimensions described above. We computed a *tag helpfulness* metric, defined as the difference between the message value for messages where the tag is true and where the tag is false. Figure 4 (right) shows the results of this analysis. Dynamics, policy, positive valence, and concrete recipes positively predict message value, while negative valence negatively predicts message value. One interpretation of this pattern is that participants who knew how to do well on the task wrote about both how the game works and what to do to succeed, while participants who did not have useful advice to give wrote about doing poorly or getting stuck.

Experiment 2

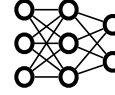
Experiment 2 was designed to validate the message value simulations. We gave participants messages from Experiment 1 and had them play each domain for 5 rounds. We were primarily interested in the extremes of the message value distribution: the messages that our simulations identified as best and worst. We selected the three highest-value messages and the two lowest-value messages for each domain and ensured that the messages came from different chains. We also added “gold” messages written by the first author to contain complete descriptions of the game rules, and a no message condition to compare behavior in the absence of an informative message.

Methods

140 participants completed the experiment on Prolific.⁴ They were paid a base payment of \$6 plus a performance-based

⁴The preregistration for this experiment can be found at <https://osf.io/5w4yx/>

Current inventory:
Tool: 🔥 stove
Tool: 💧 water
Ingredient: 🥒 cucumber, value: 0, features: vegetable
Ingredient: 🥓 raw bacon, value: 0, features: protein
Ingredient: 🌽 corn, value: 0, features: grain
Ingredient: 🌾 wheat, value: 0, features: grain
Current score: 0



```
{
  "reasoning": "I will try to cook the raw bacon
  using the stove, as cooking meat often
  increases its value and changes its features.",
  "action": ["stove", "raw bacon"]
}
```

Figure 5: Illustration of the language model simulation framework. On each timestep, the model sees the current inventory and score and produces reasoning and an action.

bonus of up to \$1 (average 47¢) and took a median of 27 minutes to complete the experiment.

Participants were given the same instructions and practice round as in Experiment 1. After completing the practice round, participants played all four domains for five rounds each. For each domain, they were randomly assigned to one of four message conditions. In the *best 3* condition, the participant read one of the messages with the three highest values. In the *worst 2* condition, the participant read one of the messages with the two lowest values. In the *gold* condition, they saw the messages written by the first author. In the *none* condition, participants simply saw “There is no message for you this round.” Example gold, best, and worst messages in the potions domain are shown in Table 1.

Results

We preregistered the hypothesis that participants would do better on average in the best 3 condition than in the worst 2 condition, measured via the main effect of condition in a Bayesian mixed-effects model.⁵ Figure 6 shows participants’ average scores by message condition. We found a significant difference between these two conditions ($\beta = 23.79$ [19.86, 27.69]). Participants also performed significantly better given gold messages than given the best 3 messages ($\beta = 20.40$ [15.40, 25.36]), and they performed slightly but significantly worse given one of the worst 2 messages compared to no message ($\beta = 5.18$ [0.04, 10.32]).

The difference between best 3 and gold conditions suggests that even the best messages were not as complete and clearly expressed as they could have been. Humans achieved notably lower average scores than the language model agents when

⁵This model had the formula `score ~ message_condition * round_num_centered + (1 + round_num_centered | domain) + (1 + round_num_centered | participant_id)`

condition	message
gold	A perfect potion has one magical ingredient and one mundane ingredient, so you should pick one of each type to combine. Before combining ingredients, you should use the vial on solid ingredients and the filter on liquid or gas ingredients. This will increase their value and turn them into liquids. Combining those liquids will get you a score of 100.
best	Solid ingredients fail being filtered. Vial + solid item. Liquid or gas item+ filter. Combine both = 100 Combine one magical with one mundane for best results
worst	Don't touch the vial...its worthless.

Table 1: Example gold, best, and worst messages for the potions domain used in Experiment 2

given good messages (gold and best 3 conditions). This may be because the language models we used in our simulations are optimized for instruction following, which might have led them to follow the instructions in the messages closely. Humans may have focused more on exploring the game environment than exploiting the knowledge in the messages.

General Discussion

The ability to discover abstract knowledge and transmit it to others is crucial to enabling the complexity of human culture, but experimental tools to study this ability have been limited. In this paper, we developed a paradigm to study the learning and transmission of abstract rules experimentally, along with a simulation framework to analyze the value of messages. Individual participants in our experiments learned many of the game rules and improved their performance over rounds, but chains of participants passing messages from one person to the next did not improve reliably. However, our simulation framework and follow-up experiment revealed that some participants discovered useful knowledge and wrote about it in a way that was helpful to recipients.

Experiment 1 was characterized by dramatic variability in participants’ performance, particularly in the chain condition. This variability could simply reflect the difficulties of online data collection, but it might also have a parallel in real-world cultural evolution. Individuals in a population vary in their skills, aptitudes, and how much effort they put into any given task. Some fields of human knowledge have been pushed forward substantially by a few people making great efforts. Perhaps learning from only one randomly-chosen teacher is not a realistic setting in which complex abstract knowledge can be discovered and refined over generations. Choices of who to learn from might be essential components of the open-ended discovery of useful abstract knowledge.

What is the relative importance of guided variation versus cultural selection? In guided variation, individual human cog-

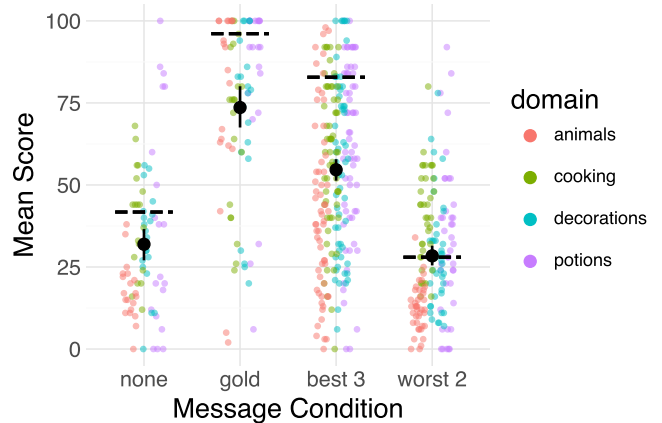


Figure 6: Participant scores by message condition in Experiment 2. Each point shows one participant’s average score in one domain and error bars show bootstrapped 95% confidence intervals. Dashed lines show mean performance of language model agents given the same messages.

nition transforms information in non-random ways; in cultural selection, some types of information are transmitted disproportionately (Mesoudi, 2021). Our crafting game may require *both* guided variation and selection for robust improvement over generations in a reasonable amount of time. Discovering useful abstract knowledge is hard, and stumbling upon it without an individual capacity to reason is highly unlikely. But even with an ability to reason, not everyone discovers many rules, so the ability to choose who to learn from might be critical for subsequent generations to improve meaningfully. This account would align with the results of prior studies which found that selecting who to learn from based on indicators of success is often an adaptive strategy, and that it allows populations to sustain more complex and effective algorithms (Laland, 2004; Thompson et al., 2022).

While the task we have developed is useful to study the transmission of abstractions, it poses difficulties in measuring participant performance. Participants appear to have different strategies for navigating the explore-exploit trade-off inherent to the game. Some participants try to learn all of the game rules in hopes of eventually getting a score of 100 every time, while others are more willing to accept lower scores. Another limitation of the current study is that there is a lot of variability by domain, but we do not understand what makes people behave differently in different domains. In particular, there is a much larger difference between immortal individual and chain conditions in the potions domain than the other domains.

In sum, here we built tools to study the cultural learning of abstract knowledge. This framework opens up many possible directions for elaboration in future work. One particularly promising direction is to test the hypothesis that selecting who to learn from in the previous generation, or combining multiple messages, will lead to more improvement over generations than the chain setup we tested here.

Acknowledgments

The authors thank the Stanford Computation and Cognition Lab and Language and Cognition Lab for feedback on early versions of this work. We extend particular thanks to Alvin Tan and Peter Zhu for giving valuable feedback on a draft of this paper. Data collection was partially supported by gift funds from Google Inc. Claude 4.7 Opus was used for code review and proofreading and Claude 4.5 Sonnet was used to assist in writing the experiment and analysis code.

References

- Agarwal, S., Ahmad, L., Ai, J., Altman, S., Applebaum, A., Arbus, E., Arora, R. K., Bai, Y., Baker, B., Bao, H., et al. (2025). Gpt-oss-120b & gpt-oss-20b model card. *arXiv preprint arXiv:2508.10925*.
- Beppu, A., & Griffiths, T. (2009). Iterated learning and the cultural ratchet. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 31(31).
- Boyd, R., Richerson, P. J., & Henrich, J. (2011). The cultural niche: Why social learning is essential for human adaptation. *Proceedings of the National Academy of Sciences*, 108(supplement_2), 10918–10925.
- Brändle, F., Binz, M., & Schulz, E. (2022). Exploration beyond bandits. In I. Cogliati Dezza, E. Schulz, & C. M. Wu (Eds.), *The drive for knowledge: The science of human information seeking* (pp. 147–168). Cambridge University Press.
- Brändle, F., Stocks, L. J., Tenenbaum, J. B., Gershman, S. J., & Schulz, E. (2023). Empowerment contributes to exploration behaviour in a creative video game. *Nature Human Behaviour*, 7(9), 1481–1489.
- Chopra, S., Tessler, M. H., & Goodman, N. D. (2019). The first crank of the cultural ratchet: Learning and transmitting concepts through language. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 41.
- Colas, C., Mills, T., Prystawski, B., Tessler, M. H., Goodman, N., Andreas, J., & Tenenbaum, J. (2025). Language and experience: A computational model of social learning in complex tasks. *arXiv preprint arXiv:2509.00074*.
- Comanici, G., Bieber, E., Schaekermann, M., Pasupat, I., Sachdeva, N., Dhillon, I., Blistein, M., Ram, O., Zhang, D., Rosen, E., et al. (2025). Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *arXiv preprint arXiv:2507.06261*.
- Derex, M., & Boyd, R. (2015). The foundations of the human cultural niche. *Nature Communications*, 6(1), 8398.
- Ham, H., Zhao, B., Griffiths, T. L., & Vélez, N. (2025). Teaching recombinable motifs through simple examples. *Cognitive Science*, 49(8), e70103.
- Harrison, P. M. C., Marjeh, R., Adolphi, F., van Rijn, P., Anglada-Tort, M., Tchernichovski, O., Larrouy-Maestri, P., & Jacoby, N. (2020). Gibbs sampling with people. *Advances in Neural Information Processing Systems*, 33.
- Laland, K. N. (2004). Social learning strategies. *Learning & Behavior*, 32(1), 4–14.
- Lewis, H. M., & Laland, K. N. (2012). Transmission fidelity is the key to the build-up of cumulative culture. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 367(1599), 2171–2180.
- Liu, A., Mei, A., Lin, B., Xue, B., Wang, B., Xu, B., Wu, B., Zhang, B., Lin, C., Dong, C., et al. (2025). Deepseek-v3.2: Pushing the frontier of open large language models. *arXiv preprint arXiv:2512.02556*.
- Martinez, J., Frank, M. C., & Haber, N. (2024). Modeling social learning through demonstration in multi-armed bandits. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 46.
- Mesoudi, A. (2016). Cultural evolution: A review of theory, findings and controversies. *Evolutionary biology*, 43(4), 481–497.
- Mesoudi, A. (2021). Cultural selection and biased transformation: Two dynamics of cultural evolution. *Philosophical Transactions of the Royal Society B*, 376(1828), 20200053.
- Morgan, T. J., & Feldman, M. W. (2025). Human culture is uniquely open-ended rather than uniquely cumulative. *Nature Human Behaviour*, 9(1), 28–42.
- Prystawski, B., Arumugam, D., & Goodman, N. (2023). Cultural reinforcement learning: A framework for modeling cumulative culture on a limited channel. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 45(45).
- Prystawski, B., Arumugam, D., & Goodman, N. D. (2025). Lossy communication constrains iterated learning. *arXiv preprint arXiv:2511.18220*.
- Smith, K., Kirby, S., & Brighton, H. (2003). Iterated learning: A framework for the emergence of language. *Artificial Life*, 9(4), 371–386.
- Sumers, T. R., Ho, M. K., Hawkins, R. D., & Griffiths, T. L. (2023). Show or tell? exploring when (and why) teaching with language outperforms demonstration. *Cognition*, 232, 105326.
- Tessler, M. H., Madeano, J., Tsividis, P. A., Harper, B., Goodman, N. D., & Tenenbaum, J. B. (2021). Learning to solve complex tasks by growing knowledge culturally across generations. *arXiv preprint arXiv:2107.13377*.
- Thompson, B., & Griffiths, T. L. (2021). Human biases limit cumulative innovation. *Proceedings of the Royal Society B*, 288(1946), 20202752.
- Thompson, B., Van Opheusden, B., Sumers, T., & Griffiths, T. (2022). Complex cognitive algorithms preserved by selective social learning in experimental populations. *Science*, 376(6588), 95–98.
- Tomasello, M. (1999). *The cultural origins of human cognition*. Harvard University Press.
- Witt, A., Toyokawa, W., Lala, K. N., Gaissmaier, W., & Wu, C. M. (2024). Humans flexibly integrate social information despite interindividual differences in reward. *Proceedings of the National Academy of Sciences*, 121(39), e2404928121.

- Yaman, A., Tian, S., & Lindström, B. (2025). Semantic knowledge guides innovation and drives cultural evolution. *arXiv preprint arXiv:2510.12837*.
- Zhao, B., Mieczkowski, E., Arumugam, D., Vélez, N., & Griffiths, T. (2025). Discovering hidden laws in innovation by recombination. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 47.
- Zhao, B., Vélez, N., & Griffiths, T. (2024). A rational model of innovation by recombination. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 46.