

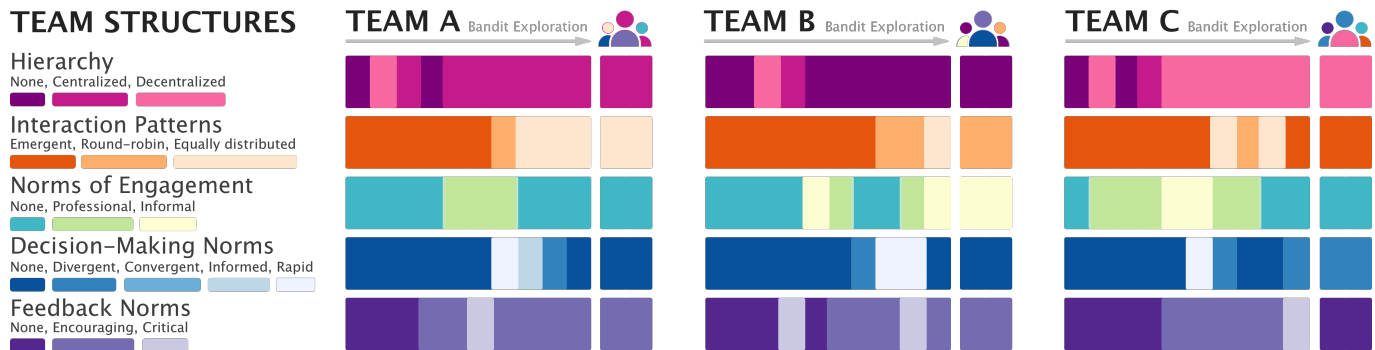


# In Search of the Dream Team: Temporally Constrained Multi-Armed Bandits for Identifying Effective Team Structures

Sharon Zhou, Melissa Valentine, Michael S. Bernstein

Stanford University

sharonz@cs.stanford.edu, mav@stanford.edu, msb@cs.stanford.edu



**Figure 1.** Each team succeeds under different roles, norms, and interaction patterns: there are no universally ideal team structures. The DreamTeam system exposes teams to a series of different team structures over time to identify effective structures for each team, based on feedback. We introduce multi-armed bandits with temporal constraints to guide this exploration without overwhelming teams in a deluge of simultaneous changes.

## ABSTRACT

Team structures—roles, norms, and interaction patterns—define how teams work. HCI researchers have theorized ideal team structures and built systems nudging teams towards them, such as those increasing turn-taking, deliberation, and knowledge distribution. However, organizational behavior research argues against the existence of universally ideal structures. Teams are diverse and excel under different structures: while one team might flourish under hierarchical leadership and a critical culture, another will flounder. In this paper, we present *DreamTeam*: a system that explores a large space of possible team structures to identify effective structures for each team based on observable feedback. To avoid overwhelming teams with too many changes, DreamTeam introduces *multi-armed bandits with temporal constraints*: an algorithm that manages the timing of exploration–exploitation trade-offs across multiple bandits simultaneously. A field experiment demonstrated that DreamTeam teams outperformed self-managing teams by 38%, manager-led teams by 46%, and teams with unconstrained bandits by 41%. This research advances computation as a powerful partner in establishing effective teamwork.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, to republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from [permissions@acm.org](mailto:permissions@acm.org).

CHI 2018, April 21–26, 2018, Montreal, QC, Canada

© 2018 Copyright held by the owner/author(s). Publication rights licensed to ACM. ISBN 978-1-4503-5620-6/18/04...\$15.00

DOI: <https://doi.org/10.1145/3173574.3173682>

## ACM Classification Keywords

H.5.3 Group and Org. Interfaces: Collaborative computing.

## Author Keywords

Teams; technical social computing; multi-armed bandits.

## INTRODUCTION

Human-computer interaction research has featured a long line of systems that influence teams' roles, norms, and interaction patterns. Roles, norms, and interaction patterns—known collectively as *team structures*—define how a team works together [32]. For many years, HCI researchers have theorized ideal team structures [1, 45] and built systems that nudge teams toward those structures, such as by increasing shared awareness [18, 20], adding channels of communication [65, 64, 70], and convening effective collaborators [38, 50]. The result is a literature that empowers ideal team structures.

However, organizational behavior research denies the existence of universally ideal team structures [53, 3, 4, 26]. Structural contingency theory [17] has demonstrated that the best team structures depend on the task, the members, and other factors. This begs the question: when should a team favor one team structure over another? Should the team have centralized or decentralized hierarchy? Should it enforce equal participation from each member? Should members offer each other more encouraging or critical feedback? The wrong decisions can doom a team to dysfunction [32, 53, 3, 4]. Even highly-paid experts—managers—struggle to pick effective team structures [15]. They are hardly to blame, as the set of possibilities is vast [29], with lengthy volumes, dedicated

handbooks, and multi-page diagrams created to tame even just one dimension of this unwieldy space [39].

In this paper, we introduce *DreamTeam*, a system that identifies effective team structures for each team by adapting teams to different structures and evaluating each fit. *DreamTeam* explores over time, experimenting with values along many dimensions of team structures such as hierarchy, interaction patterns, and norms. The system utilizes feedback, such as team performance or satisfaction, to iteratively identify the team structures that best fit each team.

Unfortunately, the state-of-the-art technical approach for this exploration results in so much simultaneous change that teams become quickly overwhelmed. Multi-armed bandits (hereafter, bandits) are a common approach for efficiently exploring different options, called arms, and exploiting the best arms over time. A network of bandits allows multiple bandits to each represent a different independent dimension (e.g., hierarchy, interaction patterns, norms for providing feedback), for which each bandit will find an optimal strategy [23, 13, 12]. The challenge is that each bandit independently explores different values of its dimension, exposing teams to changes across several dimensions at once. While theoretically optimal, this amount of change overwhelms teams, as they are not always ready to adapt rapidly [44, 41] or are only prepared to change certain dimensions at specific times [41]. *DreamTeam* requires an approach that manages a network of bandits so that both its overall change rate and dimensional change rate are sufficiently low for teams to adapt.

The core technical contribution of this paper is an algorithm for *multi-armed bandits with temporal constraints*. This algorithm models (A) along which dimensions and (B) how quickly a network of bandits can explore. The algorithm redistributes the probabilities of reward estimated with Thompson sampling [2] so that (A) the expected change within each dimension (e.g. from centralized to decentralized hierarchy) respects a constraint on *when that dimension* can change, and (B) the expected total number of changes respects a constraint on *how many* dimensions can change simultaneously. By renormalizing the probabilities from Thompson sampling, bandits with temporal constraints control the expected number of changes at each time step dimensionally and globally.

We evaluated *DreamTeam* by convening teams to complete intellectual tasks—a series of complex collaborative puzzles—across several hours. We randomized teams into five conditions: teams that chose their structures each round without instruction, collectively with instructions, with a manager, with unconstrained bandits, or with *DreamTeam* (bandits under temporal constraints). Across ten rounds, we collected teams' scores as a measure of performance. *DreamTeam* teams significantly outperformed all other conditions, by 38%–46% on average per round.

This paper contributes: (1) the concept of computationally-empowered identification of effective team structures; (2) a system manifesting this concept; (3) a network of bandits with temporal constraints, which regulates exploration timing; and

(4) an evaluation demonstrating improvements on a complex intellectual task.

## RELATED WORK

This paper draws together HCI research with organizational behavior and multi-armed bandit literature.

### Human-computer interaction and groups

Computation can convene on-demand, computationally-aided groups to achieve crisis mapping [40], interface prototyping [47], research [74, 62], writing [34, 7], sensemaking [30], and design [14, 55, 72]. These works manipulate the group's collaboration structures. Their strategies vary widely, including agile methodology [74], workflows [47], summarization [73, 36], and external idea influence [55, 72]. *DreamTeam* builds on them by acknowledging that each such structure is appropriate for some groups and goals but not others, lending an adaptive layer to these contributions.

Foundational research in HCI and CSCW demonstrated that technological mediation affects teamwork. Remote teams underperform in-person teams [45], communicate less fluidly [56], and struggle with conflict as size grows [35]. To counter these effects, researchers introduced systems to amplify the unique benefits of computer-mediated teamwork [31], such as improving group awareness [18, 20, 28] or designing alternative communication channels [65, 64]. This work has generally advocated particular team structures as ideal, e.g., that more transparency and more awareness is desirable. *DreamTeam* recognizes that the appropriate team structures vary by team and task, and proposes an approach that allows each team to find the right structures for the job.

Convening appropriate collaborators is a structure of particular importance. Team dating [38] exposes participants to each other before deciding who to work with. Like with *DreamTeam*, participants' post-dating selections depend on specific factors of the people and task. Other approaches match team members based on prior familiarity [50]. Automated approaches, however, prompt criticism and dispute over how and why the system has grouped certain people [33]. *DreamTeam* focuses instead on what happens after the team is convened: how do members identify effective collaboration structures?

### Organizational behavior: structural contingency theory

An early aim of organizational research was to understand why organizations are structured in particular ways [57, 27, 60], and researchers quickly found that different organizational structures were more or less effective under different conditions. For example, formalized vertical structures are efficient in fairly stable environments, but fail in tumultuous environments, where organizations with more organic, emergent structures are better able to adapt [37]. Organizational structures are contingent on many factors including size, scale, technology, geography, national or cultural differences, scope, individual predispositions, resource dependency, and organizational life cycles [25, 54]. This perspective is often referred to as *structural contingency theory* [37, 11]. The intuition behind structural contingency theory extends to team-level analysis as well. For example, structures promoting autonomy are useful

when teams' process interdependence is low [58], and structures that align with team members' values are most likely to be effective [67]. Team structures include roles, specialization, and hierarchy [10, 32, 68]. They are contingent on team size, task complexity, national or cultural differences, member preferences, and many other factors [66, 59, 51]. They matter for team performance because they influence information sharing, recognition of expertise and responsibilities, effectiveness of decision-making processes, and levels of conflict [10, 46, 9].

DreamTeam is designed around the recognition that team structures are malleable. Intervening once is neither sufficient nor the limit [53]: teams respond better to interventions during a task than before they assemble and begin [22]. Adapting team structures can improve team performance [41], but the timing of these changes is critical. Different structures are malleable at different times [41, 53]: some (a) when the team first forms its identity, in the first half of the task; others (b) when the team focuses on performance, under the deadline in the second half; and still others (c) when the team is working through interpersonal dynamics throughout the entire duration. Teams are vulnerable when they develop maladaptive processes during these phases, and lock themselves into poor strategies [32, 49, 5]. DreamTeam draws on this literature to determine when the system allows a given dimension to change.

### Multi-armed bandits

Multi-armed bandit algorithms are used to explore a wide array of options and identify the best one, comparable to A/B testing. Algorithms for multiple simultaneous bandits—known as a network of bandits—examine the strategy that each bandit produces or the differences between independent groups, such as treatments for patient subpopulations [23]. A network of bandits may be used to find a globally optimal configuration for the network, in which the bandits are able to share information about their response and contexts, but are otherwise independent of each other [13, 12]. Following these approaches, we model each structural dimension as a bandit, constructing a network of bandits that operate cohesively together. We introduce temporal constraints to identify the best arm for each bandit (dimension), under a global strategy.

The multi-armed bandit literature examines constrained exploration with risk-averse bandits and budgeted bandits, as well as dynamic bandits for adaptive environments. These constraints so far are not temporal, but include identifying the least risky arm to select using the mean-variance metric [52, 24, 63], or an overall budget to expend over fixed or stochastic costs on the arms [71]. While dynamic bandits adapt to varying reward distributions over time [8], they do not constrain exploration. Our work is the first to add global constraints across a network of bandits, required by the realities of human teamwork. We contribute a renormalization technique to reweigh sampled values from the posterior distribution of Thompson sampling to enforce these constraints.

### DREAMTEAM

DreamTeam aids teams in identifying the structures that are most effective for them by experimenting with different structures over time on multi-armed bandits. DreamTeam takes in a

| DIMENSION<br>BANDIT                               | VALUES<br>ARMS                                                                                                                                                                                                                                                                                                             |
|---------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Hierarchy</b><br>Early constraint              | <ul style="list-style-type: none"> <li>None: no structure enforced, teams can do anything</li> <li>Centralized: elect a leader</li> <li>Decentralized: majority-led vote to determine responses</li> </ul>                                                                                                                 |
| <b>Interaction patterns</b><br>Ongoing constraint | <ul style="list-style-type: none"> <li>Emergent: allow patterns to emerge organically</li> <li>Round-robin: take turns making suggestions</li> <li>Equally distributed: post in similar quantities as teammates</li> </ul>                                                                                                 |
| <b>Norms of Engagement</b><br>Late constraint     | <ul style="list-style-type: none"> <li>None: no structure enforced, teams can do anything</li> <li>Professional: use professional language with each other</li> <li>Informal: get to know your teammates and add fun to the task</li> </ul>                                                                                |
| <b>Decision-making Norms</b><br>Late constraint   | <ul style="list-style-type: none"> <li>None: no structure enforced, teams can do anything</li> <li>Divergent: think of diverse ideas</li> <li>Convergent: generate consensus and use compromise</li> <li>Informed: make informed and thoughtful judgments</li> <li>Rapid: make decisions as quickly as possible</li> </ul> |
| <b>Feedback Norms</b><br>Ongoing constraint       | <ul style="list-style-type: none"> <li>None: no structure enforced, teams can do anything</li> <li>Encouraging: give positive encouraging comments to teammates</li> <li>Critical: critique and play devil's advocate</li> </ul>                                                                                           |

**Table 1.** DreamTeam's dimensions, values, and temporal constraints, drawn from organizational behavior literature (e.g., [32]).

set of dimensions representing team structures, such as *hierarchy*, along with values for each dimension, such as *centralized* or *decentralized*. The system reacts based on feedback from automatically collected metrics such as task performance, from self-reported metrics such as member ratings on the team's collaboration, or on a mix. These metrics are represented in a reward function. DreamTeam learns and selects what to explore next based on this reward function, honing in on what combination of values would optimize for maximum reward. This combination represents the team structures that work well for that team.

### Network of bandits

Following bandit literature on modeling several dimensions [23, 13, 12], we equip our system with multiple bandits, each representing one dimension of team structures (Table 1). We construct a network of five bandits spanning the dimensions of hierarchy, interaction patterns, norms of engagement, decision-making norms, and feedback norms. While many different dimensions and values are possible, we generated these dimensions and values based on their prominence in the literature on team structures [32, 68]. Each dimension maps onto a set of values (Table 1).

A bandit encodes the values of the dimension that it represents as arms. For example, the hierarchy dimension has values of centralized, decentralized, or none (i.e. laissez-faire hierarchy) represented by a three-armed bandit. Together, these bandits represent the space of team structures that teams are able to explore. Across five dimensions each with three to five values (arms), there exist 405 combinations and thus 405 possible team structures to which a team can adapt on DreamTeam. As the team works together over time, each bandit collects feedback from the team, such as their performance, and considers which arm to select next. The next round's arm selection leverages a technique called Thompson sampling, a bandit algorithm that (a) uses past rewards to update every arm's Bayesian probability distribution, an arm's likelihood of returning the highest reward, and then (b) samples an arm from these probability distributions for the next round [2]. Should every bandit pick the same arm as it had in the previous round, each dimension's value would remain constant and the overall

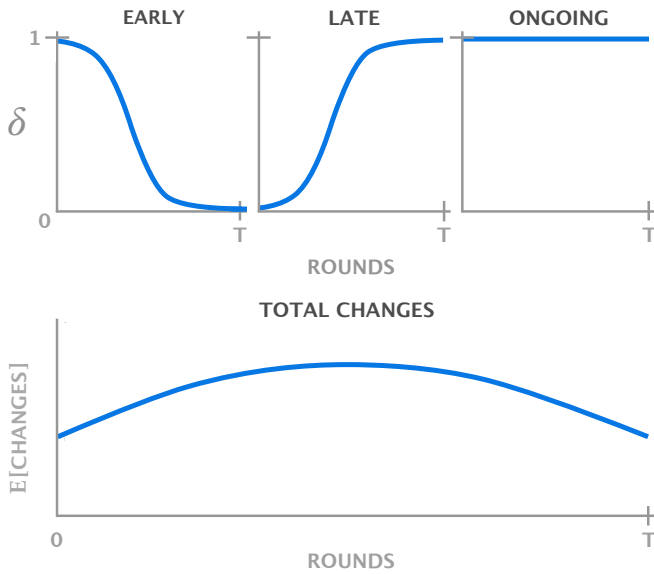


Figure 2. Temporal constraints. Top: dimensional temporal constraints apply to individual bandits, depicting the times at which each is free to explore: early, late, and ongoing. Bottom: the global temporal constraint depicts the expected allowed changes across all dimensions.

structures of the team would remain the same. Every bandit that selects a different arm from its last introduces an additional change to the structures—an explorative probe in the bandit’s dimension.

While a network of bandits using Thompson sampling results in an efficient scan over the different possible combinations, it also results in an overwhelming number of changes that teams cannot process. With each bandit individually suggesting changes, the entire network may together suggest many changes at each round—especially early on, when bandits prefer exploration over exploitation. Teams cannot adapt to so many changes to their team structures at once, nor to changes in all dimensions at once [41]. Motivated by prior work in organizational behavior [44, 41], we next construct algorithmic models that (a) pace changes by dimension and (b) pace the number of changes that bandits introduce together.

### Posterior renormalization procedure

In order to change structures without overwhelming the team, we introduce an algorithm for temporal constraints in bandit exploration. These temporal constraints reshape how many dimensions, and which dimensions, can change at each time step. Leveraging work on when teams are most receptive to changes and when they are least receptive [44], we model the total number of changes to which a team is expected to adapt, called the *global temporal constraint* (Figure 2). Based on Marks et al.’s temporality framework of when certain dimensions of team structures can adapt [41], we also model the adaptability of each dimension, which we call *dimensional temporal constraints*, by classifying each dimension to an apt time for them to change: early, late, and ongoing (Figure 2).

To model global and dimensional temporal constraints, we must develop a procedure that renormalizes the probabilities

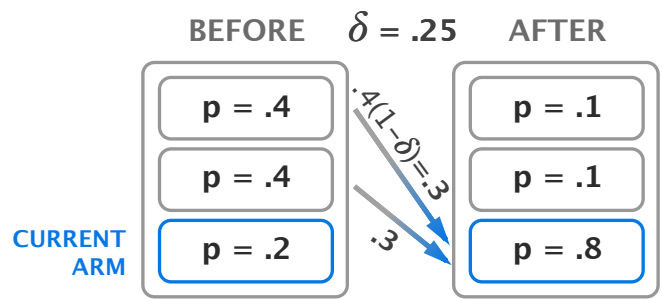


Figure 3. Posterior renormalization shifts the bandit’s probabilities to make it less likely that a dimension changes. Here  $\delta = .25$ , so the bandit becomes one-quarter as likely to explore a non-active arm.

sampled with Thompson sampling. At a high level, this procedure of posterior renormalization normalizes and shifts the sampled probabilities away from inactive arms and toward the current (selected) arm in order to lower the probability of a change. This procedure can be applied to a single dimension to control when it is allowed to change, or across all dimensions to enforce the expected number of simultaneous changes.

Thompson sampling calculates a probability distribution over the observed rewards thus far (up until time  $t$ ). It then samples  $\theta_i(t)$  from this distribution for each arm  $i$ , which represents the likelihood that arm  $i$  could be the best arm (i.e., returning the highest reward). The arm sampled with the highest believed reward payout is selected as the next arm. Once the bandit observes a reward for the current arm  $c$ , such as a score measuring the team’s performance, the observed reward rate updates the probability distribution, from which  $\theta_i(t)$  values are resampled for each arm. Selecting a different arm from  $c$  at the next time step  $t + 1$  means exploring a different value in the dimension and introducing a change to the team’s structures. To describe the procedure, we will fix the time step to  $t$  and represent the sampled value for arm  $i$  as  $\theta_i$ .

Our procedure normalizes the sampled  $\theta_i$  from Thompson sampling such that the samples reflect relative probabilities, and shifts the cumulative probability of selecting inactive arms  $\{i : i \neq c\}$  onto the current active arm  $c$  by a value  $\delta \in [0, 1]$ . This value  $\delta$  represents the discounted probability of selecting an inactive arm that remains on those inactive arms, consistent with discount factors in utility discounting for Markov decision processes, of which bandits are a subset. In other words, it is the fraction of the expected value of changes from  $c$  that is needed to achieve the temporal constraint. If there is no constraint, then no shifting is necessary,  $\delta = 0$ , and the bandit explores without modification to Thompson sampling. If  $\delta = 1$ , then all probabilities of selecting an inactive arm are shifted to  $c$ , resulting in  $\theta_c = 1$  while  $\theta_i = 0, i \neq c$ , which restricts exploration entirely.

As in Figure 3, take the example where three quarters of exploration is to be constrained ( $\delta = 0.25$ ) because behavioral research says not to tinker with this dimension—say, hierarchy—late in a team’s lifetime. Suppose the distribution of Thompson sampling probabilities was  $\theta_1 = 0.4, \theta_2 = 0.4, \theta_3 = 0.2$  and the current arm is  $c = 3$ . The procedure shifts  $1 - \delta = 0.75$  of the exploration originally on the other arms  $\theta_1$  and  $\theta_2$ , such

that  $\theta'_1 = \delta\theta_1 = 0.1$  and  $\theta'_2 = \delta\theta_2 = 0.1$ . The amount that was shifted away from them was moved onto the current arm  $\theta'_3 = \theta_3 + (1 - \delta)\theta_1 + (1 - \delta)\theta_2 = 0.2 + 0.3 + 0.3 = 0.8$ . Thus, the probability of exploring inactive arms is constrained to a quarter of its original amount while exploitation of the current arm is increased to compensate.

We detail the formal equation below. Let  $c$  specify the current arm in play,  $\theta_i$  represent the sampled probabilities on arm  $i$  of the bandit after Thompson sampling, and  $\delta$  be the shift amount. Below,  $\theta'_i$  is the resulting probability of selecting each arm from the procedure:

$$\theta'_i = \begin{cases} \theta_i \delta, & \text{if } i \neq c \\ \theta_i + \sum_{j \neq c} \theta_j (1 - \delta), & \text{if } i = c \end{cases}$$

We discuss below how we arrive at  $\delta$  for global and dimensional temporal constraints.

### Dimensional temporal constraints

*Dimensional temporal constraints* model when a given dimension is amenable to change. Following Marks et al.'s temporality framework of team structures [41], we map each dimension onto one of three opportunities to change it: early (hierarchy), late (interaction patterns, decision-making norms), and ongoing (norms of engagement, feedback norms). For example, teams are ready to adapt to hierarchical changes like determining whether they need a leader early on, but the same changes become disruptive to their ability to collaborate later.

In order to model dimensional temporal constraints, DreamTeam considers the time step  $t$  at which the system receives feedback from the team, e.g. a score of their performance, in relation to the overall time horizon  $T$ . If  $t/T$  is small, the team is early in its process, and early dimensions will be more able to change, but late dimensions should not explore as freely, and vice versa if  $t/T$  is close to 1.

We construct the following algorithmic models for each dimensional constraint (Figure 2), in which  $T$  is the time horizon of the task,  $t$  is the current time step, and  $\delta$  is the fractional expected value of changing the arm from the current one:

1. *Early*: Teams are most prepared for changes to these dimensions early along their progress on  $T$  as they begin to form their identity, becoming less adaptable to them as time progresses [41]. For example, teams are resilient to exploring different hierarchical structures early on, but less resilient to changing them later. We model early dimensions as a sigmoid, that has nearly unconstrained exploration early on ( $\delta$  is high, near 1) and that becomes constrained ( $\delta$  is low, near 0) after the team is halfway through their progress, i.e.  $\delta = 1 / (1 + e^{t-T/2})$
2. *Late*: Teams are most prepared for changes to these dimensions later in their progress on  $T$ , at first less receptive to them and becoming more so over time, as they focus on performance over forming a team identity [41]. We model late dimensions as a sigmoid similar to early dimensions', as a reflection across the midpoint parallel to the y-axis, i.e.  $\delta = 1 / (1 + e^{T/2-t})$

3. *Ongoing*: Teams are prepared for changes in these dimensions throughout the duration of the task. Such dimensions engage with interpersonal dynamics that can help teams at all times [41]. We model ongoing dimensions without constraints on exploration with  $\delta = 1$ .

This process produces a value for  $\delta$ , allowing us to now leverage posterior renormalization to restrict exploration. With these  $\delta$  values, the model adjusts the probabilities of each arm to fit the constraints, using the posterior renormalization procedure above. This is done for every bandit, thus constraining exploration for each dimension based on time. The resulting values from all dimensions are then renormalized across the bandit network.

### Global temporal constraint

While dimensional temporal constraints dictate when an individual dimension (bandit) should be changing, this does not address the global problem of too many bandits changing at once. For this global constraint, we will use the posterior renormalization procedure again across bandits.

In order to model a global constraint on all dimensions, we consider the expected value of the total number of arm changes at a given time step, and constrain that value to restrict overall exploration. We choose to use an expected value framework instead of capping the maximum because this affords bandits the opportunity to explore several extremely good arms even if these arms number above the desired constraint. We draw on prior work [44] to model the progression of a team's adaptability that suggests teams are open to more changes closer to the midpoint of their work. We thus model the global constraint as a downward-facing parabola from 0 to the time horizon  $T$ , with its highest value at  $T/2$ , the midpoint of the team's progression. The mathematical model is as follows, where  $y$  is the expected number of allowed changes (Figure 2):

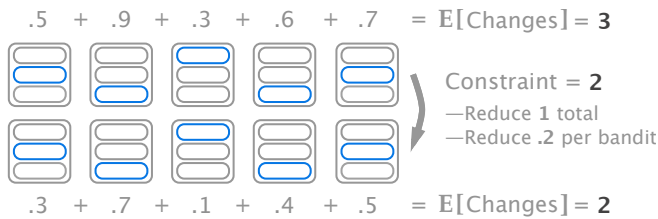
$$y \propto (t - T/2)^2$$

Using the probabilities adjusted to fit dimensional constraints, the model calculates the expected number of changes that are allowed,  $y$ , from the above equation, as well as the expected number of changes that the bandits anticipate having together  $z$ . If the number of allowed changes  $y$  is greater than the number of those anticipated  $z$ , we do not force the algorithm to explore more: this would force teams to change to arms that do not work for them and would render Thompson sampling ineffective.

On the other hand, should the anticipated number of changes  $z$  exceed the expected number allowed  $y$ , the model will shift probabilities such that the expected value is equal to the desired threshold. The model takes the excess amount  $z - y$  and distributes the burden of reducing this amount across all dimensions. The posterior renormalization procedure then distributes this reduction within each bandit, preserving any relative probabilities from Thompson sampling and dimensional constraints. The equation for finding  $\delta$  on each bandit is as follows, where  $D$  is the number of dimensions and  $z_d$  is the expected value of change for the given dimension  $d$ :

$$\delta = 1 - \frac{z - y}{z_d D}$$





**Figure 4.** Posterior renormalization on the global temporal constraint reduces the expected number of changes across the network of bandits.

Note that  $z_d$  represents the probability that a dimension will change, and is calculated by adding together the probabilities of the inactive arms  $z_d = \sum_{i \neq c} \theta_i$ , while  $z$  represents the global expected number of changes, calculated by totaling the expected values of change on all dimensions  $z = \sum z_d$ .

As illustrated in Figure 4, consider the case where the global expected number of changes is 3 across DreamTeam’s five bandits, but the allowed expected value is currently 2. The excess is 1, so each bandit needs to constrain its expected number of changes by reducing the probabilities of selecting inactive arms by  $1 / 5 = .2$ . In order to distribute this reduction proportionally across each bandit’s arms, we calculate  $\delta$  for each bandit using the equation above and renormalize the probabilities to fit the constraint.

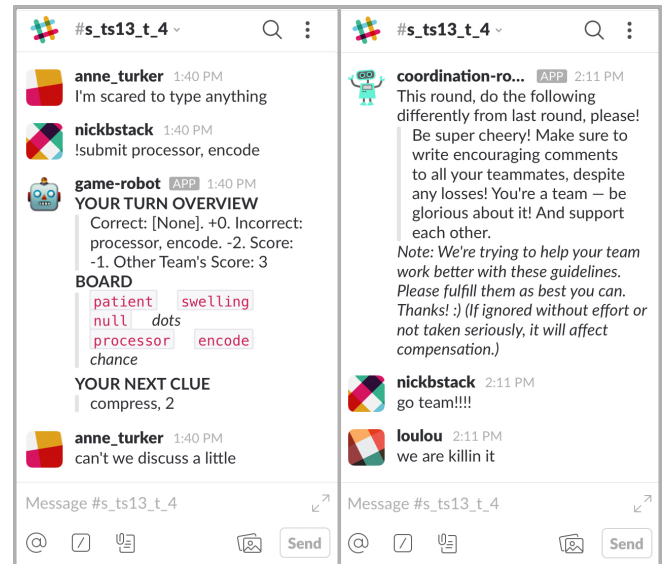
The model thus constrains exploration when the expected number of changes overreaches the amount desired. Applying to both dimensional and global temporal constraints, the core of our technical approach is to redistribute the probabilities from Thompson sampling such that they meet the temporal constraints. After accounting for both sets of temporal constraints, the system proceeds to sample arms based on their new probabilities.

### Integration

We integrate DreamTeam with Slack ([www.slack.com](http://www.slack.com)), a chat platform, using a Slack bot named coordination-robot. Coordination-robot is the user interface of DreamTeam. It joins teams in a Slack channel and offers guidance to their team structures. As team members post messages and complete tasks together, coordination-robot draws feedback from the channel’s posts, automatically taking in salient features from the team—such as their performance on a task or self-reported satisfaction. For example, after a team completes a task and their score is posted to the channel, coordination-robot reads this score as feedback.

Using temporally constrained bandits, coordination-robot decides which changes in team structures to introduce to the team at any given point. For instance, after receiving a low score on the team’s performance, coordination-robot might adapt the team’s hierarchical structure, or after receiving a high score, coordination-robot might still adapt hierarchy in an effort to help the team explore better viable options.

Coordination-robot communicates with the team through messages in the channel (Figure 5). For example, if coordination-robot were to decentralize a team’s hierarchy, it would post



**Figure 5.** Team experience within the Slack interface. Left: member *nickbstack* submits an answer to the system. Meanwhile, the user *anne\_turker* does not feel psychologically safe. Right: coordination-robot changes the team to adopt an encouraging feedback norm structure in a message to the channel. Upon completion, *anne\_turker* tells the team (not pictured): “You guys were the best team ever! Thank you.” This particular team experienced the greatest improvement in performance in our study.

to the channel: “You’re a democracy. Vote on what to submit and respect the majority vote.”

### EVALUATION

We evaluated our system by convening teams of workers, recruited on Amazon Mechanical Turk (AMT), following recent social computing systems work on teams [50, 38, 40]. We used Slack to develop a platform for collaborative conversation and integrated tasks, again following prior work [50]. These design decisions made it easier to modulate, monitor, and measure teams’ progress.

### Task design

We designed a collaborative intellectual [43] task to evaluate our system. According to McGrath’s classification of group tasks, intellectual tasks are cognitive activities focused on solving problems with a correct answer, spanning traditional science tasks in which a correct answer exists, is corroborated by data, or is supported by a jury of experts, as in a peer reviewed journal [43]. We evaluated DreamTeam in the domain of intellectual tasks to demonstrate increased performance for collaboration on complex intellectual assignments.

To design our task, we first piloted several tasks from other studies [50], including creating comic strip lines, Facebook ads, or marketing headlines. These tasks allowed us to observe upvotes or click-through rates from online audiences, but we sought a more rapid and robust feedback cycle, in which feedback would not be delayed and for which we did not have to identify appropriate feedback timing. Other considered tasks included puzzles (e.g., crosswords), but these were not selected because common answers were available online, and

because teammates frequently switched contexts away from the chat, reducing the collaborative nature of the task.

We adapted the popular board game *Codenames* for the Slack interface. The game includes a set of clue words, each of which corresponds to a group of words shown together on a game board. A team works together to determine which words on the board refer to clues that they are given. Each clue word also comes with a number. This number tallies the exact number of words on the game board that the clue refers to. In our adaptation, teams did not play competitively with each other, instead collaborating to maximize their team score.

We generated the clues and boards automatically with Empath [21], which uses vector space models to identify similar words to an input word. We took clue words from a random word generator and inputted them into Empath to produce the corresponding board words. We pre-tested all boards to ensure roughly equivalent difficulty.

### Method

We recruited 135 workers from AMT and randomized each worker into one of five conditions based on how the team structures were chosen: *control*, *collectively-chosen*, *manager-chosen*, *bandit-chosen*, and *DreamTeam-chosen*. We arranged participants into teams of three. We compensated workers \$12-\$23 because the task took 1–2 hours (<http://guidelines.wearedynamo.org>), and awarded a \$1 bonus for each round performed above the average. Teams first had two practice rounds to learn the game with no team structures imposed. Teams then engaged in ten rounds of the task, each of which had a time limit of 15 minutes.

We organized workers into teams following methodology from prior literature [50, 40, 42], inviting workers to a staging area until there are enough workers to separate them into teams of three who shared the same condition. All workers had previously received AMT qualifications for completing an individual version of the task.

Teams worked together in their dedicated channels on Slack. They used the interface to interact with their teammates, to submit answers, and to adapt to changes in their team structures. The system gave teams a board as a formatted list of words. Next, teams engaged in a 3-step cycle: (1) the system gave the team a clue, (2) the team submitted an answer, e.g. “!submit processor, encode” in Figure 5, and (3) the system showed the team their score on that answer. Each round of the game included four such cycles. Thus, across ten evaluation rounds, teams experienced 40 cycles. At the end of each round, the system gave the team their final performance. At the start of a new evaluation round, coordination-robot posted the relevant changes to the team’s structures.

Teams were asked to complete two training rounds of increasing difficulty prior to the evaluation rounds, to review the instructions and learn the Slack interface. The first training round was individualized for learning the game and interface. The second training round, “Round 0”, was in a team, with all dimensions set to “None” or “Emergent”. Round 0 was used as a covariate in our analysis and as a baseline input to the

network of bandits for DreamTeam and bandit-led teams. After concluding the practice rounds, the system posted the first evaluation round of the task, including the board and the initial clue. As evaluation rounds progressed, the design differed by condition:

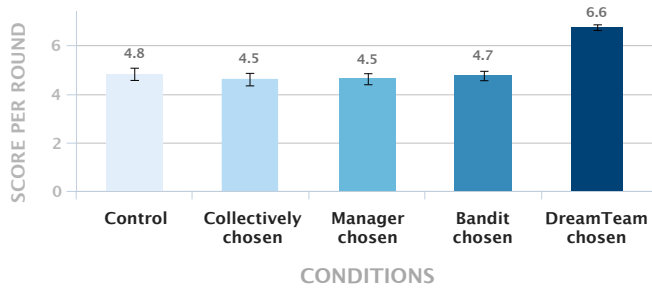
1. *Control*: No explicit structures were given. Coordination-robot did not appear.
2. *Collectively-chosen*: Teams received a list of every message that coordination-robot could use to change the team structures (i.e., every value across all dimensions). The team was instructed to post any number of these messages (none, one, or many) at the beginning of each round and, if so, these would need to be followed for that round. All team members could see the list and post. Coordination-robot did not appear.
3. *Manager-chosen*: One member of the team became the manager whose instructions the team had to follow at the beginning of each subsequent round. The manager alone had access the list of possible structural changes that coordination-robot could make. Only the manager could post a change. Managers were additionally incentivized to receive an additional bonus of \$1 for every two rounds that their team’s scores exceeded the average. Coordination-robot did not appear.
4. *Bandit-chosen*: Coordination-robot posted changes, if any, to the start of each round, powered by a network of bandits with Thompson sampling, without temporal constraints.
5. *DreamTeam-chosen*: Identical to the bandit-chosen condition, except for temporal constraints were imposed on the bandit algorithm powering coordination-robot.

For DreamTeam and bandit-chosen teams, coordination-robot was the front-end interface. The network of bandits decided what coordination-robot would post next, and evaluated the team and its current structures based on their scores on the task. At the end of the final round (ten) of all conditions, team members answered whether they would work with each other again.

### Results

Thirty-five teams completed the task with seven teams per condition. We measured the average performance of teams on all ten evaluation rounds by condition (Figure 6). DreamTeam-chosen teams ( $\mu = 6.6$ ,  $\sigma = 1.3$ ) outperformed control teams by 38% ( $\mu = 4.8$ ,  $\sigma = 2.1$ ), outperformed manager-chosen by 46% ( $\mu = 4.5$ ,  $\sigma = 2.0$ ), outperformed collectively-chosen teams by 45% ( $\mu = 4.5$ ,  $\sigma = 2.2$ ), and outperformed bandit-chosen teams by 41% ( $\mu = 4.7$ ,  $\sigma = 2.2$ ).

To analyze the distinction in performance across conditions, we conducted a repeated measures ANCOVA. We used each team’s final training round (non-intervention) score as a covariate, to control for a team’s initial performance before they experienced any interventions. There was a significant main effect of performance across conditions:  $F(4, 40) = 5.66$ ,  $p < 0.01$ . We performed post hoc Tukey tests to examine pairwise differences between conditions. DreamTeam significantly outperformed all other conditions (all  $p < .05$ ). No other conditions



**Figure 6.** DreamTeam-chosen teams had higher performance than teams in any other condition on a complex intellectual task: ANCOVA  $p < 0.05$ , all pairwise comparisons to Dreamteam-chosen  $p < 0.05$ .  $N = 45$ .

were significantly different from each other. Taken together, these results indicate that DreamTeam teams outperformed all other conditions, including bandit-chosen teams that had no temporal constraints.

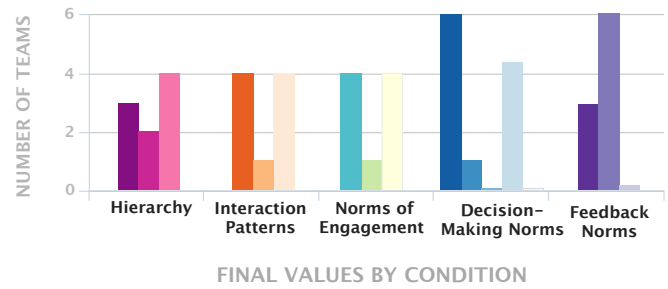
#### *Teams' most effective structures varied substantially*

What can we learn about the team structures that were most appropriate for this task? We inspected the final team structures for teams in the DreamTeam-chosen condition (Figure 7). *No two teams had the same final team structures.* These results lend support to structural contingency theory: the right set of team structures depends on each team, even for teams working on the exact same task. Strikingly, for each dimension, an average of 4.2 teams (of nine total) diverged from the most frequently chosen value. We evaluated each dimension's variation by taking the probability distributions of their values and calculating their entropy, where 1 means a uniform distribution with no leading candidate, and 0 means a completely predictable distribution with one option most effective for every team. The dimensions' entropy measurements: hierarchy (.97), interaction patterns (.88), norms of engagement (.88), decision-making norms (.53), and feedback norms (.58). Overall, dimensions had either moderate or high entropy, and most had high entropy—suggesting high variation across teams along those dimensions.

The dimensions with lower entropy suggest some patterns work well for this task. Of the six teams, none had a final configuration with a critical feedback norm, and most had an encouraging feedback norm. Only one team was effective with round-robin interaction patterns, while the others were divided evenly into equally distributed and emergent interaction patterns. Most teams wound up with no explicit decision-making norm: just one team was effective under an informed norm, two teams converging on a divergent norm, and no teams were found to be effective with convergent or rapid decision-making norms. Along the dimension of hierarchy, teams were fairly scattered: decentralized (4), centralized (2), and none (3).

#### *Teams adhered to altered structures unless overwhelmed*

Teams on the whole followed and listened to coordination-robot without complaint, taking the changes seriously. All DreamTeam teams adhered to the suggestions despite no real-time enforcement. However, about half of the teams in the bandit-chosen condition began to ignore coordination-robot over time. Without temporal constraints, bandit-chosen teams



**Figure 7.** The final team structures varied across DreamTeam teams: no values were effective for all teams. Colors refer to specific values shown in Table 1 and Figure 1.

had too much to absorb at once, and eventually teams lost trust in the suggested changes. In fact, the strongest performing teams in the bandit-chosen condition ultimately ignored the suggestions, even with the same wording of the messages.

Teams did not equally welcome all suggested structures—even structures that were the most effective for those teams. The encouraging feedback norm in particular was met with some derision, with some teams jokingly overemphasizing their positivity. However, while some members may not have taken the norm seriously at first, the norm had positive effects for other teammates, who increased their participation afterwards. Most teams in the DreamTeam-chosen condition ultimately converged on the encouraging feedback norm.

DreamTeam has limited insight into each team's existing dynamics and no principled understanding of each team structure, so occasionally it would make recommendations that were redundant and in effect, explore needlessly. For instance, the algorithm would suggest an encouraging feedback norm, but the team was already practicing an encouraging feedback norm that emerged implicitly. These occurrences were rare, provided the large space of options, but could decrease trust in the algorithm. It is also possible that the system might choose the wrong arm for the sake of exploration. This could actively hurt the team and cause them to engage in unpredictable conflicts. While we did not observe this, we note that there is a possibility, for example, that a team with a hostile environment could be told to change their team structure to include a critical feedback norm.

#### *Would you work with your team again?*

Upon completion of the game, we asked teams if they would work with the same team again, and that we would use this data to match them together in the future. The majority of teams across conditions said that they would work with their teammates again. While every team in the DreamTeam-chosen condition was eager to collaborate again on their team, a minority of teams across the other conditions did not echo the same enthusiasm. Some expressed neutrality, and others spouted unpleasant comments about their teammates.

#### *Conditions' strengths and weaknesses*

Control teams varied in performance. A handful of teams naturally collaborated or grew to collaborate. However, in many instances, there was little dialogue or effort to collaborate.



Collectively-chosen teams also had variation in how teams organized themselves. One team had a member rise up to a leadership position and delegate roles, assuming the task of organizing information around halfway through the team's progress. In most, however, despite the nominal collective decision on which team structures to use, there was no coordination and lack of organization. Some teammates exhibited and justified unchecked chaotic behavior, which became the norm. For instance, in response to "he submits them before i can read them", another teammates retorted "it works well in WoW [World of Warcraft], gotta spam those keys".

Manager-chosen teams would demonstrate relatively higher performance levels when the manager chose team structures that the team agreed to include, but most managers were hesitant to explore, using at most three different structural combinations and choosing one, or in a couple cases two, team structures at a time. Managers generally appeared to be loss-averse: in one team, the manager explored two options across two rounds, stuck with the one that showed a marginally better score, and decided to stop exploring for the rest of the game. The team ultimately had no discussion on the game, submitting in a round-robin. There was variation in management styles: while some managers relied on the team for feedback, asking questions like "Do you like the first way better?", and went with the popular vote, others decided firmly without asking for feedback and observing the team as they went. Some managers even chose and maintained certain team structures that their teammates did not like.

Bandit-chosen teams demonstrated polarizing performance levels. Some teams did very poorly, among the poorest across all teams. Yet, others exhibited performance levels similar to teams in DreamTeam, but these teams had begun ignoring the changes halfway throughout the task. For example, when coordination-robot changed the team's structure to have centralized hierarchy on one team, one teammate asked "who want[s] to be the leader?" to which no one replied as the team continued submitting answers and the game pressed forward.

The most striking difference between DreamTeam-chosen teams from those in other conditions was the consistency with which they synchronized, mobilized collectively, and engaged with the exploration of novel structures. Several structures, such as round-robin participation or electing a leader, requires collective involvement from all members of the team. Although unconstrained bandits exposed bandit-chosen teams to a greater number of structural configurations (9.89 on average, as opposed to 8.56 on DreamTeam), these teams ignored the suggestions more than a third of the time. Often, only a single member was willing to partake in the structural change, and could not galvanize others to engage in a collective effort. As a result, DreamTeam teams explored more structures by over 30%. Comparably, all manager-chosen teams experienced minimal explicit exploration (2.04 on average), because managers did not expose their teams to more than three configurations.

### Limitations

DreamTeam-chosen teams exhibited higher performance than any other condition, and all other conditions were indistinguishable from the control. How generalizable is this result?

One limitation is that we evaluated our system on an intellectual task, and did not investigate other task categories within McGrath's framework of group tasks, such as generating tasks that involve making a plan or brainstorming ideas together [43]. We cannot yet generalize beyond the collaborative intellectual-style task that we examined.

While our task could be adapted for real world applications, this task is inherently not one that existing teams typically tackle. Moreover, some tasks (e.g., creative tasks) may have high variance in performance. If variance is high, DreamTeam would still succeed, but take longer to converge. Other tasks may have rewards that are either impossible to measure or are made known over very long time periods. In these cases, we envision that a reward function could be driven by feedback from team members or from peer teams.

Our task took most teams one to two hours. A few hours is consistent with prior work and appropriate for many crowd-based teams (e.g., [40, 50, 38]), but many teams in organizations work together for months; we cannot generalize far beyond this short-term task. Manager-chosen teams, for example, may need more time for the manager to find the right team structures, and would eventually match or outperform DreamTeam.

Additionally, while we measured performance as teams' scores on the task, we did not examine how other factors, particularly qualitative ones like self-reported team satisfaction, or combinations of factors, could impact the reward function of the bandits and take them in a different direction. These features may introduce additional variation; for example, self-reported satisfaction may vary as different team members have different opinions, making it more difficult for bandits to identify team structures that give high reported team satisfaction.

While we convened virtual teams for our evaluation, we did not examine in-person teams or hybrid virtual/in-person teams. We also recruited workers on AMT who could accomplish the qualification task on their own, controlling for a threshold of ability and understanding of the task. Finally, for the manager-chosen condition, we did not recruit professional managers. While this is more similar to prior work in crowd teams (e.g., [47]), comparing against professional managers will be a clear avenue for future work.

The dimensions and values that we chose to adapt do not necessarily suit all teams, nor encompass all possible dimensions of team structures. As we expand the decision space, it will take longer for DreamTeam to identify the most effective structures. Furthermore, while DreamTeam searches for a resulting set of team structures that is effective, the right approach may be to instead determine an effective *policy*—the process and path along which they should navigate and explore the space.

Finally, there are limits to the causal claims that we can make from our experiment. Control teams have their team structures emerge implicitly, but the act of having teams decide by a collective effort, by an individual manager, or by a bot can itself affect how teams view their team identity. Nevertheless, we chose these comparisons, because they are the current modus operandi of self-managing and manager-led teams today.

## DISCUSSION

DreamTeam teams succeeded because they consistently took exploration to heart and engaged with the experimental process. Organization behavior literature corroborates the final converged values. For instance, one team had a member who felt psychologically unsafe to contribute, but the system introduced an encouraging feedback norm, drastically improving the member's participation [19]. Another had fairly equal participation at the outset, fitting the system's ultimate recommendation of equally distributed interaction patterns [16]. A third team had trouble engaging in discussion, but the introduction of informal norms of engagement reduced tensions around debate [61].

In contrast, bandit-chosen teams were deluged with information and as a result, ignored coordination-robot's suggestions to help them coordinate their joint efforts [6]. The variation in performance among the other conditions also reveals that some teams found success, but that several teams instead became bound to dysfunctional structures, failing to explore further. This was in spite of the fact that bandit-chosen, manager-chosen, and collectively-chosen teams explicitly received more structural options. This is consistent with the literature, in that teams will often fall into poor structures that doom them to continual failure if left uninterrupted, because the inertia to maintain the status quo overpowers the effort to initiate change [32, 49, 5].

DreamTeam teams had substantially different team structures from one another, yet still outperformed teams in other conditions. This result reinforces structural contingency theory: the most effective team structures depend on the members composing the team and how they interact. But how do we converge on those rapidly or make our best first guess? Technical approaches such as contextual bandits, which blend generalization capabilities from machine learning classifiers with multi-armed bandits, may allow the system to learn the answer to this question over time, allowing DreamTeam to identify which team structures are worth exploring first.

A core theoretical tenet of structural contingency theory is that groups must continuously evolve their structures to react to changes in the environment [17]. As described here, DreamTeam is focused on identifying an initially appropriate set of structures. Even as it gains confidence, Thompson sampling on bandits nevertheless retains some probability of experimenting with other arms, so over time DreamTeam would identify if the team needs to evolve. Exogenous events like new team members might prompt the system or team to re-trigger exploration. However, we suggest that it would be more powerful to give users greater control over resetting the system after a major exogenous event. This would allow the system to retain some agency in exploring and avoid pigeonholing, but keep users in control of some of its behaviors.

We observe three main avenues for future work. First, while we followed methodologies from prior work on studying short-term deployments in organizational behavior experiments [69] and crowd teams for organization [50], we also hope to examine teams in the real world on their actual tasks that might consume a larger time horizon. To do this, we can expand

DreamTeam to represent a set of team structures that matches the variation seen within the organization. Specifically, we would like to examine which structures real organizations are interested in exploring, over structures that they would like to remain fixed based on their corporate culture, e.g. many organizations have instituted a centralized hierarchy. We aim to extend the system to partner with managers in exploring and identifying what collaborative structures work well for their teams. DreamTeam represents a step towards computationally augmenting human managers and self-managed teams.

Second, we hope to run DreamTeam on teams that have either identified themselves or have been externally identified as dysfunctional to observe whether we can unfreeze their current environment and adapt them to a more effective team environment. An open question is whether DreamTeam will be more effective for teams that are moderately functional—and are thus more open to experimentation—or for teams that are known to be dysfunctional.

Third, the set of team structures included in the system is far from complete. Our goal was to demonstrate the approach for a set of team structures that the organizational behavior literature identified as particularly salient. However, we must now consider the broader space of team structures. A rich literature of handbooks and literature reviews provides a vocabulary for this larger set of dimensions and values (e.g., [39]). With a larger space to explore, the DreamTeam system will take longer to identify an effective structure. It may be possible to guide the system toward identifying some dimensions as principal components first, and then iterate from there.

## CONCLUSION

Effective teamwork is a wicked problem [48]: it cannot be planned in advance, and requires adaptation to the people, task, and environment. Prior work focused on pre-selecting effective structures, such as convening the right members or using the right collaborative system. Acknowledging that this correct identification is impossible in the limit, we introduce a system that helps teams adaptively identify a set of roles, norms, and interaction patterns that is effective for them. To enable DreamTeam, we contribute a model for multi-armed bandits with temporal constraints, which ensures that these structures evolve at a feasible pace and at the right period in a team's overall arc. Evidence from our evaluation suggests that teams using DreamTeam can be far more effective than existing modes of organizing for these rapidly convened teams.

We envision computation as a partner in helping groups achieve their goals. It can aid us when we exhibit biases or limited self-knowledge—such as identifying effective team structures—and it will help us re-plan when the environment shifts. DreamTeam represents a step toward this future of computationally augmented teams and organizations.

## ACKNOWLEDGMENTS

We thank the Mechanical Turk workers for their participation. This work was supported by the National Science Foundation (IIS-1351131), the Office of Naval Research (N00014-16-1-2894), DARPA, and Stanford MediaX.

## REFERENCES

1. Mark S Ackerman. 2000. The intellectual challenge of CSCW: the gap between social requirements and technical feasibility. *Human-computer interaction* 15, 2 (2000), 179–203.
2. Shipra Agrawal and Navin Goyal. 2012. Analysis of Thompson sampling for the multi-armed bandit problem. In *Conference on Learning Theory*. 39–1.
3. Deborah G Ancona, Gerardo A Okhuysen, and Leslie A Perlow. 2001. Taking time to integrate temporal research. *Academy of Management Review* 26, 4 (2001), 512–529.
4. Holly Arrow. 1997. Stability, bistability, and instability in small group influence patterns. *Journal of Personality and Social Psychology* 72, 1 (1997), 75.
5. Sasha A Barab and Thomas Duffy. 2000. From practice fields to communities of practice. *Theoretical foundations of learning environments* 1, 1 (2000), 25–55.
6. Jeremy B Bernerth, H Jack Walker, and Stanley G Harris. 2011. Change fatigue: Development and initial validation of a new measure. *Work & Stress* 25, 4 (2011), 321–337.
7. Michael S. Bernstein, Greg Little, Robert C. Miller, Björn Hartmann, Mark S. Ackerman, David R. Karger, David Crowell, and Katrina Panovich. 2015. Soylent: A Word Processor with a Crowd Inside. *Commun. ACM* 58, 8 (July 2015), 85–94.
8. Omar Besbes, Yonatan Gur, and Assaf Zeevi. 2015. Non-stationary stochastic optimization. *Operations research* 63, 5 (2015), 1227–1244.
9. Peter Michael Blau. 1974. *On the nature of organizations*. John Wiley & Sons.
10. J Stuart Bunderson and Peter Boumgarden. 2010. Structure and learning in self-managed teams: Why “bureaucratic” teams can be better learners. *Organization Science* 21, 3 (2010), 609–624.
11. Tom E Burns and George Macpherson Stalker. 1961. The management of innovation. (1961).
12. Stéphane Caron and Smriti Bhagat. 2013. Mixing bandits: A recipe for improved cold-start recommendations in a social network. In *Proceedings of the 7th Workshop on Social Network Mining and Analysis*. ACM.
13. Nicolo Cesa-Bianchi, Claudio Gentile, and Giovanni Zappella. 2013. A gang of bandits. In *Advances in Neural Information Processing Systems*. 737–745.
14. Joel Chan, Steven Dang, and Steven P Dow. 2016. Improving crowd innovation with expert facilitation. In *Proceedings of the 19th ACM Conference on Computer-Supported Cooperative Work & Social Computing*. ACM, 1223–1235.
15. Hans De Bruijn and Ernst et al. Ten Heuvelhof. 2010. *Process management: why project management fails in complex decision making processes*. Springer Science & Business Media.
16. C. K. W. De Dreu and M. A. West. 2001. Minority dissent and team innovation: The importance of participation in decision making. *Journal of Applied Psychology* 86, 6 (2001), 1191–1201.
17. Lex Donaldson. 1999. *The normal science of structural contingency theory*. SAGE Publications Ltd. 51–70 pages.
18. Paul Dourish and Victoria Bellotti. 1992. Awareness and coordination in shared workspaces. In *Proceedings of the 1992 ACM conference on Computer-supported cooperative work*. ACM, 107–114.
19. A. Edmondson. 1999. Psychological safety and learning behavior in work teams. *Administrative Science Quarterly* 44, 2 (1999), 350–383.
20. Thomas Erickson and Wendy A Kellogg. 2000. Social translucence: an approach to designing systems that support social processes. *ACM transactions on computer-human interaction (TOCHI)* 7, 1 (2000), 59–83.
21. Ethan Fast, Binbin Chen, and Michael S Bernstein. 2016. Empath: Understanding topic signals in large-scale text. In *Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems*. ACM, 4647–4657.
22. Colin M Fisher. 2017. An ounce of prevention or a pound of cure? Two experiments on in-process interventions in decision-making groups. *Organizational Behavior and Human Decision Processes* 138 (2017), 59–73.
23. Victor Gabillon, Mohammad Ghavamzadeh, Alessandro Lazaric, and Sébastien Bubeck. 2011. Multi-bandit best arm identification. In *Advances in Neural Information Processing Systems*. 2222–2230.
24. Nicolas Galichet, Michele Sebag, and Olivier Teytaud. 2013. Exploration vs exploitation vs safety: Risk-aware multi-armed bandits. In *Asian Conference on Machine Learning*. 245–260.
25. R.T. Golembiewski. 2000. *Handbook of Organizational Behavior, Second Edition, Revised and Expanded*. Taylor and Francis.
26. Google. 2016. re:Work with Google. <https://rework.withgoogle.com/print/guides/5721312655835136/>. (2016).
27. Alvin Ward Gouldner. 1959. *Organizational analysis*. Basic Books.
28. Carl Gutwin and Saul Greenberg. 2002. A descriptive framework of workspace awareness for real-time groupware. *Computer Supported Cooperative Work (CSCW)* 11, 3 (2002), 411–446.
29. J Richard Hackman. 2012. From causes to conditions in group research. *Journal of Organizational Behavior* 33, 3 (2012), 428–444.
30. Nathan Hahn, Joseph Chang, Ji Eun Kim, and Aniket Kittur. 2016. The Knowledge Accelerator: Big picture thinking in small pieces. In *Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems*. ACM, 2258–2270.

31. Jim Hollan and Scott Stornetta. 1992. Beyond being there. In *Proceedings of the SIGCHI conference on Human factors in computing systems*. ACM, 119–125.
32. Daniel R Ilgen, John R Hollenbeck, Michael Johnson, and Dustin Jundt. 2005. Teams in organizations: From input-process-output models to IMO models. *Annu. Rev. Psychol.* 56 (2005), 517–543.
33. Farnaz Jahanbakhsh, Wai-Tat Fu, Karrie Karahalios, Darko Marinov, and Brian Bailey. 2017. You Want Me to Work with Who?: Stakeholder Perceptions of Automated Team Formation in Project-based Courses. In *Proceedings of the 2017 CHI Conference on Human Factors in Computing Systems*. ACM, 3201–3212.
34. Aniket Kittur, Boris Smus, Susheel Khamkar, and Robert E Kraut. 2011. Crowdforge: Crowdsourcing complex work. In *Proceedings of the 24th annual ACM symposium on User interface software and technology*. ACM, 43–52.
35. Aniket Kittur, Bongwon Suh, Bryan A Pendleton, and Ed H Chi. 2007. He says, she says: conflict and coordination in Wikipedia. In *Proceedings of the SIGCHI conference on Human factors in computing systems*. ACM, 453–462.
36. Travis Kriplean, Michael Toomim, Jonathan Morgan, Alan Borning, and Andrew Ko. 2012. Is this what you meant?: promoting listening on the web with reflect. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. ACM, 1559–1568.
37. Paul R Lawrence and Jay W Lorsch. 1967. Differentiation and integration in complex organizations. *Administrative science quarterly* (1967), 1–47.
38. Ioanna Lykourantzou, Robert E. Kraut, and Steven P. Dow. 2017. Team Dating Leads to Better Online Ad Hoc Collaborations. In *Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing (CSCW '17)*. ACM, New York, NY, USA, 2330–2343.
39. Thomas W Malone, Kevin Crowston, and George Arthur Herman. 2003. *Organizing business knowledge: The MIT process handbook*. MIT press.
40. Andrew Mao, Winter Mason, Siddharth Suri, and Duncan J Watts. 2016. An experimental study of team size and performance on a complex task. *PloS one* 11, 4 (2016).
41. Michelle A Marks, John E Mathieu, and Stephen J Zaccaro. 2001. A temporally based framework and taxonomy of team processes. *Academy of management review* 26, 3 (2001), 356–376.
42. Winter Mason and Siddharth Suri. 2012. Conducting behavioral research on Amazon's Mechanical Turk. *Behavior research methods* 44, 1 (2012), 1–23.
43. Joseph Edward McGrath. 1984. *Groups: Interaction and performance*. Vol. 14. Prentice-Hall Englewood Cliffs, NJ.
44. Gerardo A Okhuysen and Mary J Waller. 2002. Focusing on midpoint transitions: An analysis of boundary conditions. *Academy of Management Journal* 45, 5 (2002), 1056–1065.
45. Gary M Olson and Judith S Olson. 2000. Distance matters. *Human-computer interaction* 15, 2 (2000), 139–178.
46. Derek Salman Pugh and David J Hickson. 1976. *Organizational structure in its context*. Saxon House Westmead, Farnborough.
47. Daniela Retelny, Sébastien Robaszkiewicz, Alexandra To, Walter S Lasecki, Jay Patel, Negar Rahmati, Tulsee Doshi, Melissa Valentine, and Michael S Bernstein. 2014. Expert crowdsourcing with flash teams. In *Proceedings of the 27th annual ACM symposium on User interface software and technology*. ACM, 75–85.
48. Horst WJ Rittel and Melvin M Webber. 1973. Dilemmas in a general theory of planning. *Policy sciences* 4, 2 (1973), 155–169.
49. Sandra L Robinson and Anne M O'Leary-Kelly. 1998. Monkey see, monkey do: The influence of work groups on the antisocial behavior of employees. *Academy of Management Journal* 41, 6 (1998), 658–672.
50. Niloufar Salehi, Andrew McCabe, Melissa Valentine, and Michael Bernstein. 2017. Huddler: Convening Stable and Familiar Crowd Teams Despite Unpredictable Availability. In *Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing (CSCW '17)*. ACM, New York, NY, USA, 1700–1713.
51. Jane E Salk and Mary Yoko Brannen. 2000. National culture, networks, and individual influence in a multinational management team. *Academy of Management journal* 43, 2 (2000), 191–202.
52. Amir Sani, Alessandro Lazaric, and Rémi Munos. 2012. Risk-aversion in multi-armed bandits. In *Advances in Neural Information Processing Systems*. 3275–3283.
53. Michaëla C Schippers, Amy C Edmondson, and Michael A West. 2014. Team reflexivity as an antidote to team information-processing failures. *Small Group Research* 45, 6 (2014), 731–769.
54. W.R. Scott and G.F. Davis. 2015. *Organizations and Organizing: Rational, Natural and Open Systems Perspectives*. Taylor & Francis.
55. Pao Siangliulue, Joel Chan, Steven P Dow, and Krzysztof Z Gajos. 2016. IdeaHound: Improving Large-scale Collaborative Ideation with Crowd-powered Real-time Semantic Modeling. In *Proceedings of the 29th Annual Symposium on User Interface Software and Technology*. ACM, 609–624.

56. Jane Siegel, Vitaly Dubrovsky, Sara Kiesler, and Timothy W McGuire. 1986. Group processes in computer-mediated communication. *Organizational behavior and human decision processes* 37, 2 (1986), 157–187.
57. HA Simon. 1947. *Administrative behavior; a study of decision-making processes in administrative organization*. Macmillan.
58. Christine A Sprigg, Paul R Jackson, and Sharon K Parker. 2000. Production teamworking: The importance of interdependence and autonomy for employee strain and satisfaction. *Human Relations* 53, 11 (2000), 1519–1543.
59. Bradley R Staats, Katherine L Milkman, and Craig R Fox. 2012. The team scaling fallacy: Underestimating the declining efficiency of larger teams. *Organizational Behavior and Human Decision Processes* 118, 2 (2012), 132–142.
60. James D Thompson. 1967. *Organizations in action: Social science bases of administrative theory*. Transaction publishers.
61. H. H. M. Tse, M. T. Dasborough, and N. M. Ashkanasy. 2008. A multi-level analysis of team climate and interpersonal exchange relationships at work. *Leadership Quarterly* 19, 2 (2008), 195–211.
62. Rajan Vaish, Snehal Kumar Neil S Gaikwad, Geza Kovacs, Andreas Veit, Ranjay Krishna, Imanol Arrieta Ibarra, Camelia Simoiu, Michael Wilber, Serge Belongie, Sharad Goel, James Davis, and Michael S Bernstein. 2017. Crowd Research: Open and Scalable University Laboratories. In *Proceedings of the 32nd annual ACM Symposium on User Interface Software and Technology*.
63. Sattar Vakili and Qing Zhao. 2016. Risk-Averse Multi-Armed Bandit Problems under Mean-Variance Measure. *IEEE Journal of Selected Topics in Signal Processing* 10, 6 (2016), 1093–1111.
64. Gina Venolia, John Tang, Ruy Cervantes, Sara Bly, George Robertson, Bongshin Lee, and Kori Inkpen. 2010. Embodied social proxy: mediating interpersonal connection in hub-and-satellite teams. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. ACM, 1049–1058.
65. Fernanda B Viégas and Judith S Donath. 1999. Chat circles. In *Proceedings of the SIGCHI conference on Human Factors in Computing Systems*. ACM, 9–16.
66. Ruth Wageman. 1995. Interdependence and group effectiveness. *Administrative science quarterly* (1995), 145–180.
67. Ruth Wageman and Frederick M Gordon. 2005. As the twig is bent: How group values shape emergent task interdependence in groups. *Organization Science* 16, 6 (2005), 687–700.
68. Ruth Wageman, J Richard Hackman, and Erin Lehman. 2005. Team diagnostic survey: Development of an instrument. *The Journal of Applied Behavioral Science* 41, 4 (2005), 373–398.
69. Lan Wang, Jian Han, Colin M Fisher, and Yan Pan. 2017. Learning to share: Exploring temporality in shared leadership and team learning. *Small Group Research* 48, 2 (2017), 165–189.
70. Terry Winograd. 1986. A language/action perspective on the design of cooperative work. In *Proceedings of the 1986 ACM conference on Computer-supported cooperative work*. ACM, 203–220.
71. Yingce Xia, Haifang Li, Tao Qin, Nenghai Yu, and Tie-Yan Liu. 2015. Thompson Sampling for Budgeted Multi-Armed Bandits. In *IJCAI*. 3960–3966.
72. Lixiu Yu, Aniket Kittur, and Robert E Kraut. 2014. Distributed analogical idea generation: inventing with crowds. In *Proceedings of the 32nd annual ACM conference on Human factors in computing systems*. ACM, 1245–1254.
73. Amy X Zhang, Lea Verou, and David R Karger. 2017b. Wikum: Bridging Discussion Forums and Wikis Using Recursive Summarization. In *CSCW*. 2082–2096.
74. Haoqi Zhang, Matthew W Easterday, Elizabeth M Gerber, Daniel Rees Lewis, and Leesha Maliakal. 2017a. Agile Research Studios: Orchestrating Communities of Practice to Advance Research Training. In *Companion of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing*. ACM, 45–48.