



Understanding Human Dynamic Sampling Objectives to Enable Robot-assisted Scientific Decision Making

SHIPENG LIU, Department of Electrical and Computer Engineering, University of Southern California, USA
CRISTINA G. WILSON, Collaborative Robotics and Intelligent Systems Institute, Oregon State University, USA

BHASKAR KRISHNAMACHARI, Department of Electrical and Computer Engineering, University of Southern California, USA

FEIFEI QIAN*, Department of Electrical and Computer Engineering, University of Southern California, USA

Truly collaborative scientific field data collection between human scientists and autonomous robot systems requires a shared understanding of the search objectives and tradeoffs faced when making decisions. Therefore, critical to developing intelligent robots to aid human experts, is an understanding of how scientists make such decisions and how they adapt their data collection strategies when presented with new information *in situ*. In this study we examined the dynamic data collection decisions of 108 expert geoscience researchers using a simulated field scenario. Human data collection behaviors suggested two distinct objectives: an information-based objective to maximize information coverage, and a discrepancy-based objective to maximize hypothesis verification. We developed a highly-simplified quantitative decision model that allows the robot to predict potential human data collection locations based on the two observed human data collection objectives. Predictions from the simple model revealed a transition from information-based to discrepancy-based objective as the level of information increased. The findings will allow robotic teammates to connect experts' dynamic science objectives with the adaptation of their sampling behaviors, and in the long term, enable the development of more cognitively-compatible robotic field assistants.

Additional Key Words and Phrases: decision making, robot-assisted scientific exploration, human cognitive model

1 INTRODUCTION

Robots and rovers are beginning to assist human field scientists in a range of planetary and earth science mission tasks [21, 30]. Mobile robotic platforms can move through increasingly complex natural environments while measuring environment properties, bringing the precision of laboratory experimentation to the field [4, 13, 24–26, 31]. Due to the uncertainty and dynamic nature of some natural environments, multi-stream data need to be considered simultaneously to understand the mechanisms underlying natural processes. The use of robots for data collection has already helped further geoscientists' understanding of desert sediment dynamics. For example, in prior work, we deployed a legged robot, RHex [28], to assist human geoscientists in field data collection across deserts in NM and CA [25], to reveal how environmental properties such as soil strength and vegetation density influence sediment transport and desertification processes. The robot provided human geoscientists

*Corresponding author, feifeiqi@usc.edu

Authors' addresses: Shipeng Liu, shipengl@usc.edu, Department of Electrical and Computer Engineering, University of Southern California, Los Angeles, California, USA; Cristina G. Wilson, wilsoncr@oregonstate.edu, Collaborative Robotics and Intelligent Systems Institute, Oregon State University, Corvallis, OR, USA; Bhaskar Krishnamachari, bkrishna@usc.edu, Department of Electrical and Computer Engineering, University of Southern California, Los Angeles, CA, USA; Feifei Qian, feifeiqi@usc.edu, Department of Electrical and Computer Engineering, University of Southern California, Los Angeles, California, USA.

Permission to make digital or hard copies of part or all 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 third-party components of this work must be honored. For all other uses, contact the owner/author(s).

© 2023 Copyright held by the owner/author(s).

2573-9522/2023/9-ART

<https://doi.org/10.1145/3623383>

with high spatiotemporal resolution data on leg-soil interactions *in-situ*. Access to this data allowed scientists to test hypotheses about soil strength in the field, update their beliefs, and dynamically adapt their data collection strategies to enable important scientific discoveries [35].

Most state-of-the-art robots assisting in data collection – including the robot RHex used in our research [25, 35] – are used as mobile sensor suites, taking low-level command from humans to execute the navigation, sensing and sampling, while human experts bear the full burden of integrating and interpreting data for future data collection decision making. Cognitive science literature has shown that when this burden exceeds the processing capacity of the human mind, experts are likely to rely on mental shortcuts and rules-of-thumb (heuristics) [5, 20] and have trouble flexibly adapting thinking and behavior in response to new information [2]. This leaves expert scientists vulnerable to decision biases that can lead to missed scientific discoveries [34]. Robots could help human teammates to make better decisions by taking on increased responsibility in collaborative exploration, e.g., through autonomous low-level data collection decisions or in-situ decision support.

Here we take the position that, in order for robots to truly understand and anticipate expert data collection needs and recognize potential human decision pitfalls, we must first understand how scientists make and adapt data collection decisions. In this study, we combine human cognitive experiments and quantitative robot reward functions, to allow robots to connect the adaptation of scientists’ sampling behaviors with their dynamic science objectives. We study the cognition and behavior of a diverse pool of possible expert geoscientist end-users, to understand how experts update data collection objectives in response to incoming information. The observed human abstract objectives are then represented as hypothesized reward functions for robots to predict corresponding sampling locations. These hypothesized reward functions serve as a simplified and testable model of human behavior, and allow systematic investigation of the effect of different decision parameters. Comparison between robot-predicted sampling locations and human-experts’ decision data reveals key decision factors governing experts’ data collection strategies, and informs how experts tradeoff different objectives to select and adapt strategies in response to new data. With this work, we take the first necessary step towards a robot inferring scientist objectives and predicting data collection behavior; where the aim is not to replicate how humans make decisions, but rather for a robot to have a representation of thought processes (theory of mind) that lead to scientists’ actions or goals [7]. In this manner, robots can begin to move from mobile sensor suites to more supportive and intelligent teammates – with the ability to infer and flexibly support experts’ dynamic objectives, or even identify biased behavior and provide targeted support.

2 BACKGROUND

The majority of past research has approached mobile robotics for science exploration and data collection as a coverage (mapping) problem [32] or a simple search (target detection) problem [27]. Much progress has been made in these areas with autonomous path planning solutions, for example using adaptive sampling algorithms [11, 14] to generate trajectories that maximize information gain [10], minimize uncertainty [17, 33], or minimize risk [12, 23] while minimizing costs. Some research has integrated expert prior knowledge with adaptive sampling algorithms; for example, by having scientists represent their prior beliefs as spatial probabilities on a “hypothesis map” [3] that can be used to generate trajectories that maximize scientific information gain [22]. Existing adaptive sampling algorithms present implicit tradeoffs between various sampling objectives – e.g., exploration-exploitation [9], shortest-smoothest-safest [16] – but we do not yet understand how such tradeoffs align with those made by human scientists during field data collection. In the absence of this understanding, robots cannot predict how algorithmic tradeoffs in sampling objectives might differ from human tradeoffs and therefore cannot support human teammates when their judgments are vulnerable to bias.

The primary approach taken to align algorithmic tradeoffs in sampling objectives with how humans dynamically prioritize objectives has been to learn from human input. While methods such as coactive learning [31] and

inverse reinforcement learning [1, 15] could be used to infer human priorities over different trajectories, the goal of these methods often focus on finding a reward function that allows the robot to best imitate human demonstrations. Without a general representation of decision process, the selected optimization could sensitively depend on the scientists' preference from the training set, making it challenging to cope with variations from a broader set of human subjects. Methods from human modelling studies [8] have shown promise in formulating the representation of thought process, and equipping robots with explicit models of why and how a human teammate would make one tradeoff over another. The benefit of such an approach is that the human model can directly come from human decision data, resulting in better alignment with human reasoning patterns and improved explainability [19]. The challenge, however, is how to connect the implicit human decision making principles [35] with explicit mathematical expressions that can guide robot-aided information search behaviors. As a first step towards connecting human theory of mind with robotics algorithms, the goal of our study is to find a simple, explainable principle that can help understand why humans make their sampling decisions, and how these sampling decisions were related to their scientific beliefs. Such understanding can allow future robots to understand the reasons behind human actions and offer explainable decision suggestions.

3 STUDYING HUMAN SPATIOTEMPORAL DATA COLLECTION BEHAVIOR WITH A FIELD SAMPLING INSPIRED SIMULATED SCENARIO

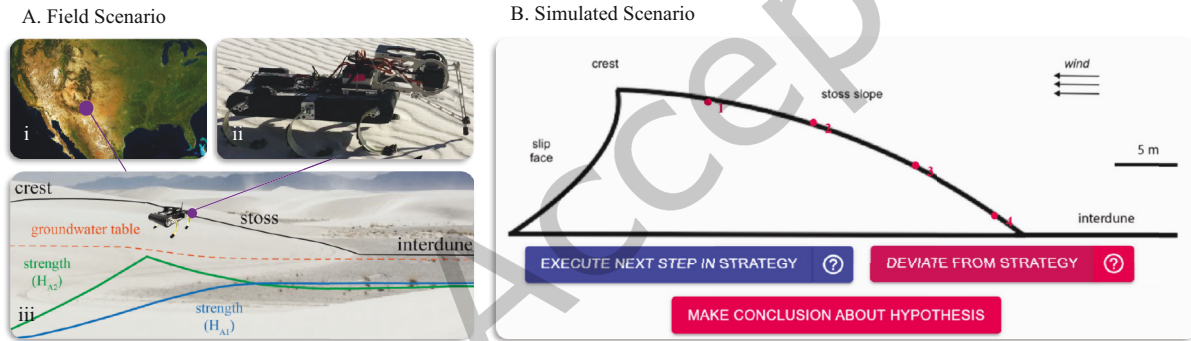


Fig. 1. (A) Field site at White Sands, NM, a dune field in the southwest of United States (i), where the RHex [28] robot (ii) assisted human scientists by collecting soil property measurements [25, 26] along a sand dune. Black line highlights the transect of a dune where we observe the largest gradient in soil properties (iii). At the crest of the dune, where the soil was driest because of its distance from the groundwater table (orange line), soil strength was expected to be low. As moisture increased on the stoss face moving towards the interdune, strength was expected to also increase before leveling off at the point of moisture saturation. This pattern of expected results, H_{A1} , is displayed in blue and is provided to participants in the simulated scenario. Through field work, geoscientists discovered that soil strength actually increases rapidly to its maximum as soil becomes slightly wet, and then decreases slightly as soil moisture becomes more saturated nearing the interdune area just before leveling off [26]. This alternative pattern of results, H_{A2} , is displayed in green. Participants in the simulated scenario were randomly assigned to sample from data sets supporting H_{A1} or H_{A2} . (B) The interactive data collection page of the web-based decision-making scenario, which was inspired from the robot-assisted field data collection scenario. Expert geoscientist participants select data collection locations on dune cross-section and measurements are provided in real-time.

Here, we use a simulated data collection scenario to determine the objectives driving expert geoscientist data collection decisions, and how changes in hypothesis beliefs (measured by subjective confidence) alter the weighting of objectives and corresponding sampling decisions. The simulated scenario allows testing human experts' decision processes within a controlled environment, where all individuals have access to the same information about

the hypothesis and measurement data. It also allows us to recruit a larger number of expert geoscientists for participation, ensuring behavioral findings are representative of the community and not idiosyncratic to a small number of scientists.

3.1 Simulated scientific data collection task

The simulated data collection scenario is inspired from real-world field data collection [26] at White Sands National Park (Figure 1A-i), a dune field in New Mexico [25, 26]. At White Sands, the availability of *in-situ* data from a field-deployable robot [26] (Figure 1A-iii) allowed experts to discover that the exhibited dependence of soil strength on moisture was significantly different from what was hypothesized in the literature [18]. Based on the literature, it was expected that dry soil would have the lowest strength, and as soil moisture increased, strength would also increase, eventually reaching a plateau as soil moisture saturates (see H_{A1} in blue, Fig. 1A-iii). However, field measurements taken by the robot (Fig. 1A-ii) revealed that as soil became slightly wet (1-3% moisture), soil strength rapidly increased to maximal, and then slowly decreased with the increasing of moisture, eventually stabilizing to a plateau (see H_{A2} in green, Fig. 1A-iii).

Based on the real-world data collection, we designed two data sets of soil strength and moisture, supporting H_{A1} and H_{A2} , respectively. These data sets¹ were used in our simulated scenario² to generate robot measurements at participant-selected sampling locations. In our simulated scenario, the expected hypothesis, H_{A1} , was provided to all participants. But each participant was randomly assigned to draw from one of the two data sets supporting H_{A1} or H_{A2} . Participants selected an initial data collection strategy for testing the hypothesis at a single dune, by clicking a location, l , on an image of a dune cross-section (Fig. 1B) and entering the number of measurements they wished to take at the location, n . We discretized the dune transect (length scale tens of meters) into 22 sampling locations, $l \in [1, 22]$, where 1 corresponds to the crest and 22 corresponds to the end of the interdune. Participants were allowed to click anywhere on the continuous dune transect, to indicate that they would like the robot to collect soil strength and soil moisture data from this location. The web-based scenario records the position of this clicked location on the transect image, and bins it to the nearest discretized location³, l . Data were then presented one location at a time, and participants were asked to provide their confidence towards the given hypothesis following each measurement: very confident the data supports the hypothesis, moderately confident, slightly confident, I don't have enough information, not at all confident. At any time, participants could deviate from their initial strategy, or stop data collection and make a conclusion about the hypothesis.

Once a conclusion was made, participants completed a follow-up survey that included questions about decision strategies, data interpretation, and demographic information (e.g., age, gender, years of practice). For data interpretation questions, participants were shown plots of high density data (220 total measurements, 10 measurements at 22 locations along transect) either supporting H_{A1} or H_{A2} and asked to draw a conclusion: “Does the data support the previously stated hypothesis? How confident are you?” Questions about decision strategies were open-response: “Did you take a consistent number of measurements at each location? Explain why you took the number you did.” and “Did you select locations that were roughly evenly spaced? Explain why you selected that spacing you did?”. After an initial examination of all participant responses, we developed categories of stated objectives, and then multiple people independently coded a sub-set (roughly 45%) of participant responses based on these categories to ensure agreement. Indices of inter-rater reliability are $k = .847$ and $.622$ for inter-rater agreement on sample number reasoning and sampling location interval reasoning question response categories, respectively, which suggested substantial agreement according to Fleiss classification [6]. Based on the result of

¹A copy of the data sets is available online here: <https://github.com/QianLabUSC/HumanRobotExploration>.

²The simulated scenario is available online here: <http://seas.upenn.edu/~foraging/transect/> (username and password are both rhex).

³It was found that the results remained qualitatively similar with different binning resolutions.

this test, we confirmed that the coding system reliably distinguished between categories of objectives, and then a single person used the same system to code the remaining participant responses.

3.2 Expert geoscientist participants

108 expert geoscientists were recruited to participate in the simulated task, hosted online. Potential participants were recruited at the 2019 Fall meeting of the American Geophysical Union, through email listerv, or personal contact with the authors. Experts had to have obtained a bachelor degree in a geoscience-related field to qualify. Our reported human data sets were from 48 female participants, 59 male participants, and 1 participant who reported their gender identity as "other". Experts ranged in age 22-72 and in their years of experience post-bachelors 0-51.

3.3 Basic performance in the task

Among the 108 experts, 55 were assigned to sample from data supporting the provided hypothesis, H_{A1} , and 53 were assigned to sample from data supporting the alternative unknown hypothesis, H_{A2} . Among the participants who received data supporting the provided hypothesis, 49 made the correct conclusion and the remaining 6 made false alarms (type I error). While among those who received data supporting the alternative unknown hypothesis, only 31 made the correct conclusion to reject and the remaining 22 made a type II error, failing to reject the hypothesis when it was false. Since the primary goal of this study is to understand how experts' beliefs are connected with their sampling behaviors, for the analyses in Sec. 4 and Sec. 5 we were not particularly concerned with the accuracy of hypothesis conclusions so much as how experts data collection decisions helped them to reach their conclusion, and how they updated their hypothesis beliefs in response to incoming data.

4 KEY DECISION FACTORS AND REWARD FUNCTIONS GOVERNING EXPERTS' INITIAL DATA COLLECTION STRATEGY

We observed that approximately 94% of participants ($n = 102$) chose uniform location intervals (equal spacing heuristic), and 85% of participants ($n = 92$) chose a constant number of samples at each location (magic number heuristic) during initial strategy selection. Furthermore, we found that approximately 67% of the participants ($n = 62$) chose a constant number of samples, $n \in [3, 6]$, and 64% of the participants ($n = 65$) chose uniform intervals, $\Delta l \in [2, 4]$ (Fig. 2C, D). This is consistent with the observations in our previous work, where we found that experts heavily rely on these heuristics during both simulated and real-world field data collection [35].

Participants responses to follow-up survey questions about decision strategy (Fig. 2A, B) revealed the majority of experts using the magic number heuristic (71%, $n = 75$) and the equal spacing heuristic (60%, $n = 69$) reported their behavior was driven by efficiency: they either mentioned reducing time and costs associated with increasing measurements or adequately capturing underlying sample variability, or balancing both.

In combination, experts' heuristic behavior and the qualitative responses they gave for relying on heuristics, suggest their decisions are driven by reward-cost trade-offs: balancing information-gathering efficiency and sampling cost while increasing information coverage during initial data collection. Based on the observed heuristics, we propose that experts use an information measure that satisfies the following features: (i) exhibits diminishing returns, where the additional information that experts expect to gain from a new sample decreases, when they continue to sample from the same location (Fig. 2E); and (ii) contains indirect information inference, where experts expect to gain information about the measurements at a location l by sampling at a nearby location, l_s ; but the expected information gained from the nearby location decreases as the distance between l and l_s increases (Fig. 2F).

The diminished information reward representation provides an explanation for the observed magic number heuristic. As the number of samples at the same location increases, the amount of new information with each

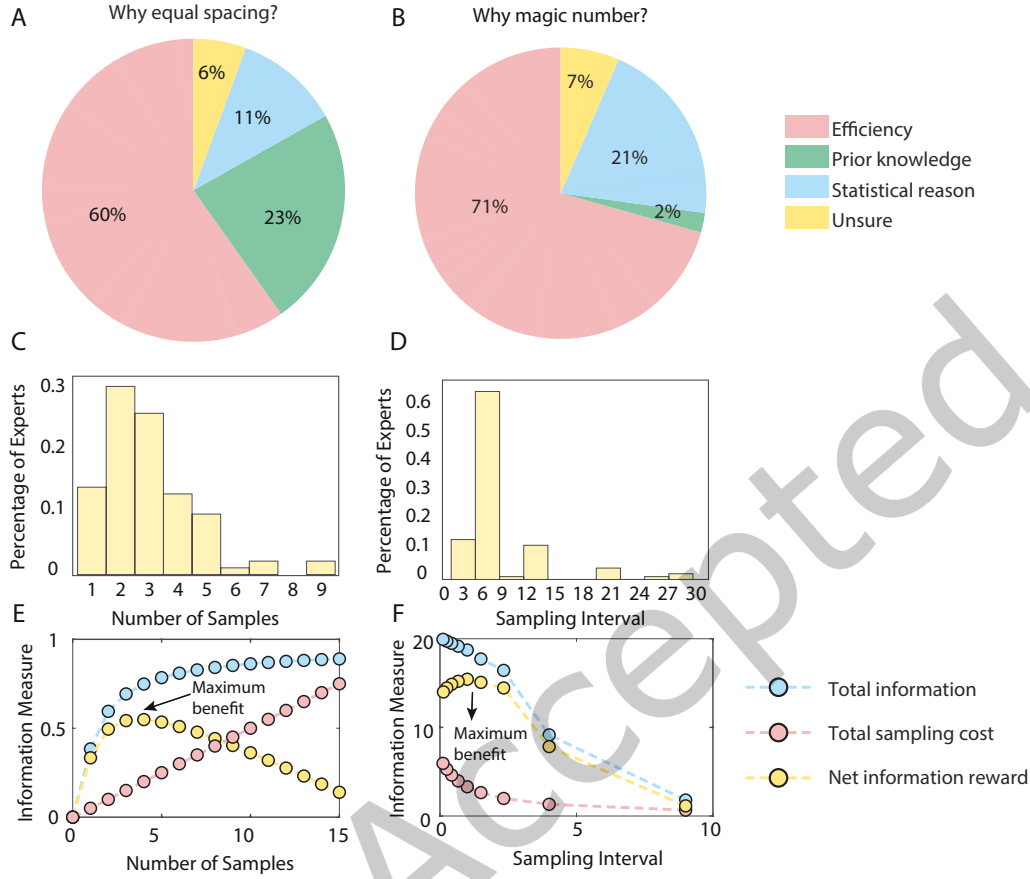


Fig. 2. (A, B) Coded qualitative responses about decision strategies from experts who took a heuristic strategy of choosing either a consistent “magic number” of samples at each location, $n = 92$, or evenly-spacing sampling locations, $n = 102$. Responses were coded as being driven by: “Efficiency” if they mentioned reducing time and costs associated with increasing measurements, and/or adequately capturing underlying sample variability; “Prior-knowledge”, if they mentioned being informed by expertise brought from their own research, and/or expectations about the hypothesis; “Statistical reason”, if they mentioned applying a particular statistical analysis; and “Unsure” if they were unable to report a reason. (C, D) Observed sampling strategy heuristics from human responses. (E) Relationship between the information measure, I , and the number of samples, n . The blue curve represents the information measure, I , at location l . The red curve represents the sampling cost to collect n samples. The yellow curve represents the net information reward, defined as the information (blue) minus the sampling cost (red). (F) Relationship between the information measure and sampling interval. The blue curve represents the information measure summed from all locations for a given sampling interval, d . The red and yellow curves represent the sampling cost and net information reward, respectively. In both (E) and (F), the sampling cost to obtain each measurement is assumed to be a constant, 0.05, across all locations. For the simple cases illustrated here, the hyperparameters were kept as constants, with their values (signal variance $\sigma_f = 10$, length scale $k = 5$) determined from the human initial sampling strategy data.

additional sample decreases (Fig. 2 E, blue, rate of change). Assuming the sampling cost for each measurement is constant, the total sampling cost to collect n samples would be proportional to n (Fig. 2 E, red). As a result, the

sample number dependent net reward (Fig. 2 E, yellow), defined as the total information (Fig. 2 E, blue) minus the total sampling cost (Fig. 2 E, red), would exhibit a non-monotonic relationship with the number of samples. The net reward (Fig. 2 E, yellow) is expected to increase the most rapidly with the first few samples. Subsequently, as the number of samples increases further, the rate of information increase slows down, and the net reward diminishes after reaching a peak.

The information inference reward representation provides an explanation for the observed equal spacing heuristic. Since information could be inferred from nearby sampled locations, as the distance between sampling locations, d , increases (*i.e.*, from densely-sampled to sparsely-sampled), the total amount of information decreases in a nonlinear fashion (Fig. 2F, blue). Assuming the sampling cost is a constant at each location, the total sampling cost (Fig. 2F, red) is inversely proportional to the sampling interval. As a result, the information inference based net reward (Fig. 2 F, yellow), defined as the total information (Fig. 2 F, blue) minus the total sampling cost (Fig. 2 F, red), would exhibit a non-monotonic relationship with the sampling interval: when sampling too densely (*i.e.*, small interval), the amount of information is sufficiently large, but the total sampling cost is large, resulting in a smaller net reward; on the other hand, when sampling too sparsely (*i.e.*, large interval), the amount of inferred information decreases rapidly, resulting in significantly-reduced total information, and consequently a smaller net reward as well. This is consistent with experts choosing a sampling location interval $d \in [2, 4]$ (Fig. 2 D) to efficiently increase coverage by inferring information about nearby locations.

Here we construct a quantitative expression of the hypothesized information measure based on the observed expert heuristics. These heuristics are neither good nor bad, rather they may lead to optimal or suboptimal outcomes depending on their fit with the decision environment. That said, an information measure that reflects humans' mental representations of information value can help us to answer the following questions: are experts' sampling behaviors primarily driven by the objective to effectively increase information? if not, what other key decision factors are driving experts' sampling decisions? This knowledge is essential if robotic teammates are to understand experts' dynamically-evolving sampling objectives and be able to suggest sampling strategies that align with experts' objectives.

In our scenario, we use the Gaussian Processes (GP) [29] to estimate the measurements, X , at different locations. Let N be a set of $|N|$ locations where we want to estimate the measurement means. The variance of estimated mean, δ^2 , is calculated as the diagonal elements of the covariance matrix of the estimated mean at those locations, $cov(X)$ (this matrix is of size $|N| \times |N|$). Let M be a set of $|M|$ locations where samples have been taken (multiple measurements at the same location will be treated as different sample points and thus M can have redundant elements). The covariance matrix of estimated mean is defined as follows:

$$cov(X) = K_{N,N} - K_{N,M} \cdot [K_{M,M} + \sigma_\epsilon^2 I]^{-1} \cdot K_{M,N} \quad (1)$$

In the above, $K_{N,M}$ denotes the $|N| \times |M|$ matrix of covariances evaluated at all pairs of sample points and estimation points. $K_{N,N}$, $K_{M,M}$ and $K_{M,N}$ are defined similarly. Each element in those covariance matrices is calculated based on squared exponential function (Radial Basis Function) as follows:

$$K(l_1, l_2) = \sigma_f^2 \cdot \exp\left(\frac{-(l_1 - l_2)^2}{2k}\right) = \sigma_f^2 \cdot \exp\left(\frac{-d^2}{2k}\right) \quad (2)$$

where k , σ_f^2 and σ_ϵ^2 are three hyperparameters, which are length-scale, signal variance, and noise variance respectively [29], and l_1 and l_2 are any two locations.

The uncertainty at a particular location l is described by the variance in the estimated mean, δ^2 . To satisfy the diminished information and information inference features, we represent the information measure as a function of δ :

$$I(l) = e^{-\delta(l)^2}. \quad (3)$$

This information measure, I , is similar to the concept of differential entropy commonly used in information theory. A notable characteristic of these two measures is that they exhibit inverse tendencies with each other, particularly in our context where they are both solely associated with the variance. However, different from the differential entropy, the proposed information measure is bounded $\in [0, 1]$, allowing us to compare information levels across different human participants. Fig. 3 illustrates the comparison between the information measure (red) and entropy (blue) for example measurements.

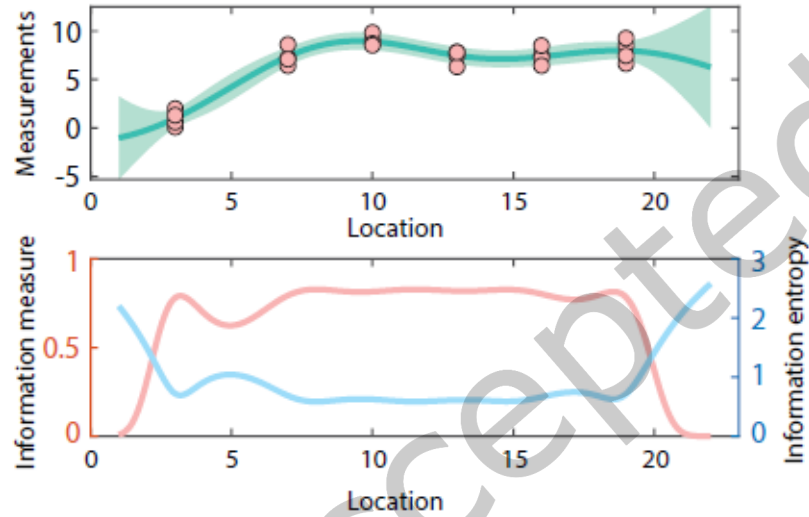


Fig. 3. A numerical example illustrating the comparison between the information measure used in this study, with the differential entropy commonly used in information theory. (A) Example measurements for illustrating the two different information measures. Red markers represent measurements at the sampling locations. The green solid curve represents the estimated mean of the measurements. The Green shaded area represents the estimated variance of the measurements. (B) A comparison between the information measure (red) and entropy (blue) for the example measurements shown in A.

By directly sampling at a location l , experts can increase the information measure at l by reducing δ^2 at that location⁴. To quantitatively illustrate this, we consider a simple case where we compute the change of information measure I at a location l , as the number of samples, n , increases at l . As shown in Fig. 2E (blue curve) and as can be shown rigorously using Gaussian Process theory, additional measurements at a given location reduce the variance in the estimated mean at that location. With the increased number of samples at the same location, the amount of reduced δ^2 (or increased information measure) with each sample decreases significantly. Considering the increased sampling cost associated with more measurements (Fig. 2E, red curve), the non-monotonic net reward (Fig. 2E, yellow curve) would result in the “magic number” heuristics observed from human experts (Fig. 2 C) to balance information coverage and sampling cost.

Similarly, it can be shown using Gaussian Process theory that additional measurements at a given location, l_s , can reduce the variance in the estimated mean at other nearby locations, l , that have significant correlation with the location where the measurements are taken. The amount of variance that can be reduced at l depends on the distance between l and l_s — the larger the distance, the smaller the reduction. This is quantitatively illustrated in

⁴We note that here δ^2 is the variance of the estimated mean, instead of the variance of the measurements themselves. The latter is an intrinsic property of the data set while the former can be reduced with the increased number of measurements.

Fig. 2 F, where we compute the total information measure from all locations, as the sampling interval increases. Due to the increased distance between locations, the total information measure decreases (Fig. 2 F, blue) with the increased sampling interval. Considering the increased sampling cost associated with shorter sampling interval (Fig. 2F, red curve), the non-monotonic net reward (Fig. 2F, yellow curve) would result in the sampling interval heuristics observed from experts' initial strategies, where the inferred information from nearby locations allows experts to gather sufficient information without obtaining measurements from each and every location.

With this information measure that reflects human's mental representation of information value, in Sec. 4.1 we construct an information based reward function, to determine whether (and if so, when) experts' sampling behavior was primarily governed by the need to increase information coverage.

4.1 Information reward function capture experts' data collection decisions at low information coverage

To determine whether experts chose sampling locations with the primary objective to effectively increase information coverage, we compare expert-selected sampling locations with the locations with largest information reward. The information reward $R(l)$ is defined as the potential increase of information level, I_{total} , if the next measurement were to be sampled from location l . The information level, $I_{total} \in [0, 1]$, was computed as the information measure summed from all locations, then normalized by the number of locations.

Here we let the robot algorithm compute the distribution of possible information reward, $R_I(l)$, among all possible sampling locations, L , and select up to 3 locations with the largest information reward peaks, L_I , to compare with experts' choices of sampling location when they adapt their initial sampling strategy. For each robot-selected location $l_{rb} \in L_I$, we compute the prediction accuracy, a_I , as a function of the distance between robot-selected sampling location, l_{rb} , and expert-selected sampling location, l_{ep} :

$$a_I = 1 - \frac{|l_{rb} - l_{ep}|}{N}, \quad (4)$$

where N is the total number of sampling locations along a transect (*i.e.*, 22 in our case). $a = 1$ corresponds to the scenario when the robot-selected location matches exactly the expert-selected location.

Fig. 4 shows the prediction accuracy of robot-selected sampling locations as compared to experts' choices from our human decision data set. If experts' sampling behaviors were primarily governed by the need to increase information, the prediction accuracy should remain high at all information levels. However, we noticed that the information based reward only captured human's sampling behaviors at low information level: prediction accuracy is much higher (an averaged accuracy of 0.91 ± 0.20) for experts who adapted their initial strategy at relatively lower information level ($I_{total} < 0.5$), whereas for experts who adapted their initial sampling strategy after they have collected a large amount of information ($I_{total} > 0.5$), the information based reward is less successful at capturing human-selected sampling locations (an averaged accuracy of 0.59 ± 0.32).

This suggested a dynamic shift in human expert sampling objectives in response to incoming data. The objective to increase information coverage seems to be driving experts' sampling strategy only during the early stage of the sampling process (*i.e.*, $I_{total} < 0.5$). Once information level increased to a sufficient amount (*i.e.*, $I_{total} > 0.5$), experts appeared to shift objectives. We hypothesized that during the later stage of sampling process, experts have formed a strong belief about the measurements, and their sampling behaviors may be governed by the objective of validating their beliefs. We investigate this hypothesis in Sec. 5.

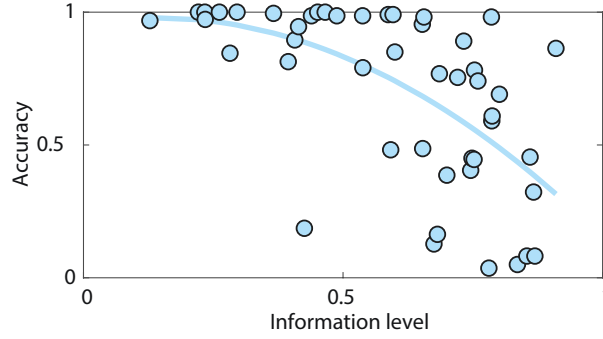


Fig. 4. Accuracy of robot-generated sampling locations, using the information reward functions, at predicting expert geoscientist data collection decisions across different information coverage levels. Each blue marker represents one expert's sampling location choices when adapt from their initial sampling strategy. Blue curve is a second order polynomial regression to provide a visual aid of the data trend.

5 KEY DECISION FACTORS AND REWARD FUNCTIONS GOVERNING EXPERTS' SAMPLING STRATEGY ADAPTATION

To understand what governs experts' sampling behavior during the later stage of sampling process, we investigated the connection between the human-selected sampling locations and their confidence towards the given hypothesis. We posited that, at high information level, expert-selected sampling locations were no longer primarily driven by the need to increase information, but were instead governed by their beliefs towards the hypothesis. To test this hypothesis, we first construct a reward function to represent experts' implicit beliefs, and then evaluate with our human data.

We asked experts to report their confidence towards the given hypothesis after each measurement. At low information level, most experts reported that they did not have enough information to make a confidence judgment (Fig. 5A, black). As the information level increased, experts began to report specific beliefs towards the given hypothesis: either that the given hypothesis (Fig. 1A iii, H_{A1}) was supported by their measurements (Fig. 5A, blue), or that it was refuted (Fig. 5A, green). We also observed that at high information levels, experts' reported beliefs towards the hypothesis were largely correlated with the amount of overall discrepancy⁵ between their collected measurements and the given hypothesis: when the overall discrepancy between the data and hypothesis was high ($> 50\%$), the majority of experts reported they were not at all confident in the hypothesis (Fig. 5B, green), and when discrepancy was low ($< 50\%$) the majority of experts reported they were slightly, moderately, or very confident in the hypothesis (Fig. 5B, blue). This suggested that the discrepancy between hypothesis and measurement may be used as a quantitative measure that represent experts' confidence towards the given hypothesis.

To determine how experts' sampling location choices are influenced by their beliefs towards the given hypothesis, here we construct a quantitative expression of the potential discrepancy for sampling at a location. The expression will allow us to compute the distribution of the discrepancy based reward among possible sampling locations, and evaluate whether experts' sampling location choices are influenced by their beliefs towards the given hypothesis, as we hypothesized (Sec. 5.1). The potential discrepancy, D , at a location l , is defined as the

⁵The overall discrepancy is computed as the root mean square error (RMSE) between the experts' measurements and the given hypothesis, following analysis reported in [35]. This overall discrepancy was computed during our post analysis and was not provided to the experts at the time of sampling and conclusion.

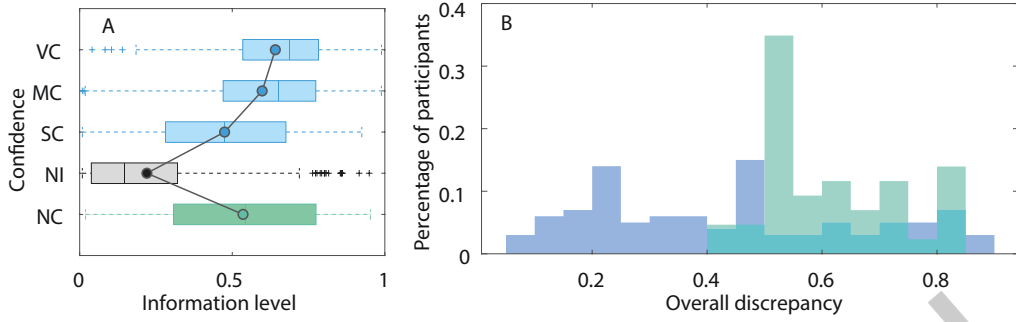


Fig. 5. Experts' reported beliefs towards the given hypothesis in response to incoming data. (A) Experts' reported confidence towards the given hypothesis across information coverage. VC, MC, SC, NI, NC stands for very confident, moderately confident, slightly confident, not enough information, and not at all confident. (B) Relationship between expert-reported beliefs towards the given hypothesis, and the overall discrepancy between the hypothesis curve, at high information level ($I \geq 0.8$). Blue represents expert-reported beliefs of data supporting the given hypothesis (reported VC, MC, or SC), whereas green represents expert-reported beliefs of data not supporting the given hypothesis (reported NC).

difference between the hypothesized soil shear strength measurement, H_y (Fig. 6B, blue dashed curve), and expected soil shear strength measurement, y (Fig. 6B, green solid curve):

$$D(I) = |H_y(I) - y(I)|. \quad (5)$$

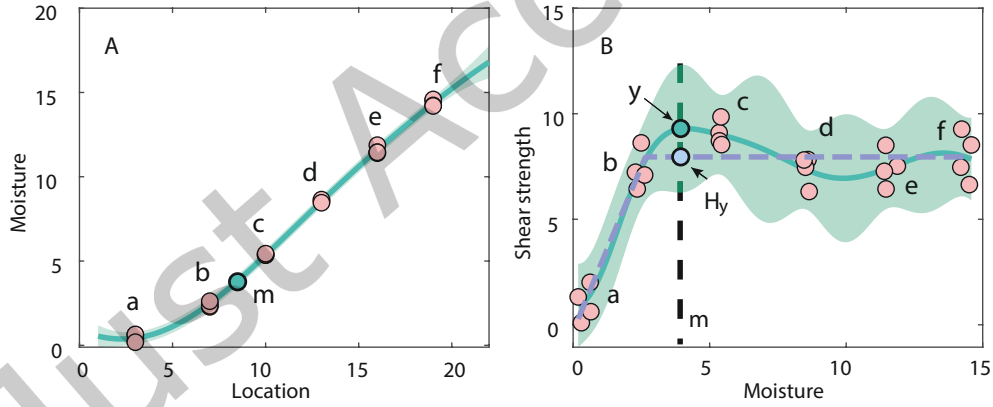


Fig. 6. Example illustration of the potential discrepancy. Red markers $a - e$ represents existing measurements. Green solid curve represents the estimated mean of the measurements. Green shaded area represents the uncertainty interval of the measurements. (A) illustrates the probability of getting a moisture measurement, m , at location l , given existing measurements $a - e$. (B) illustrates the probability of getting a soil shear strength measurement, y , at moisture m , given existing measurements $a - e$. Blue dashed curve in (B) represents H_y , the hypothesized measurements if H_{1A} was supported by the measurements. Given a location, l , the amount of difference between the solid green curve and the dashed blue curve in (B) represents the potential discrepancy, D .

To compute H_y , the given hypothesis H_{A1} (Fig. 1A iii) was represented as a piece-wise linear function of soil shear strength vs. moisture (Fig. 6B, blue dashed curve). A least-squares piece-wise linear regression was applied

for all existing measurements (Fig. 6B, *a-e*) to determine the piece-wise linear function. The expected moisture measurement, m , was modeled as a normal distribution at each location l (Fig. 6A), with its mean and variance estimated using the Gaussian Processes⁶. Similarly, the expected soil shear strength measurement, y , was modeled as a normal distribution at each moisture (Fig. 6B) using the Gaussian Processes.

We analyzed the relative discrepancy at expert-selected locations when they deviate from their initial strategy, and found an interesting shift in human’s sampling behavior: At low information level ($I_{total} < 0.5$), the amount of discrepancy at expert-selected locations exhibited no significant difference between experts who reported low confidence (Fig. 7A, red) and high confidence (Fig. 7A, blue) towards the given hypothesis. However, at high information level ($I_{total} > 0.5$), the sampling behavior bifurcated between the two belief groups. Experts who reported low confidence towards the given hypothesis selected sampling locations with significantly higher discrepancy (Fig. 7B, red, median = 0.89) as compared with experts who reported high confidence towards the given hypothesis (Fig. 7B, blue, median = 0.38).

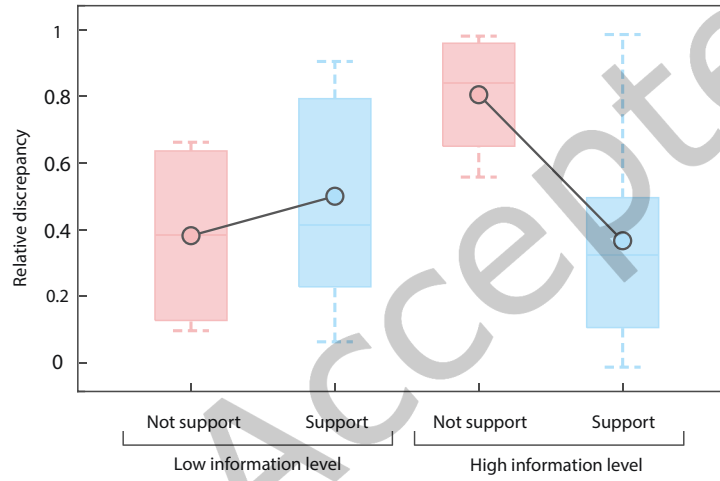


Fig. 7. Comparison of relative discrepancy between experts who reported low confidence (NC) towards the hypothesis (red boxes, “Not Support”), and experts who reported high confidence (VC, MC, SC) towards the hypothesis (blue boxes, “Support”). The relative discrepancy was computed as the potential discrepancy at expert-selected location upon sampling strategy deviation, normalized to $[0, 1]$ based on the maximal and minimal potential discrepancy across all possible sampling locations at that sampling step. The comparison was shown for both low information level ($I < 0.5$) and high information level ($I \geq 0.5$). The center mark indicates the median. The circle indicates the mean. The bottom and top edges of the box indicate the 25th and 75th percentiles, respectively. The whiskers extend to the most extreme data points not considered outliers.

This result suggested that experts’ data collection decisions at high information coverage are largely influenced by their hypothesis beliefs: We hypothesized that experts who hold the belief that their measurements do not support the given hypothesis would be more inclined to choose sampling locations with large potential discrepancy (referred to as an “invalidation” objective). In contrast, experts who hold the belief that their measurements do support the given hypothesis may choose sampling locations with smaller potential discrepancy (referred to as a “validation” objective).

Based on this result, we compute the discrepancy based reward, D , as follows: For experts holding the belief that their measurements do not support the given hypothesis, the robot uses an “invalidation” based reward function that predicts sampling locations with maximal potential discrepancy:

⁶The L-BFGS-B optimization algorithm [36] was used to determine the GP hyperparameters.

$$R_{\text{invalidate}} = \underset{l}{\operatorname{argmax}}(D(l)); \quad (6)$$

For experts holding the belief that their measurements do support the given hypothesis the robot uses a “validation” based reward function that predicts sampling locations with minimal potential discrepancy:

$$R_{\text{validate}} = \underset{l}{\operatorname{argmin}}(D(l)) \quad (7)$$

5.1 Hypothesis discrepancy reward functions captures experts’ data collection decisions at high information coverage

To evaluate the hypothesized discrepancy reward functions with our human data, we compute the distribution of discrepancy reward among all possible sampling locations. We let the robot algorithm select up to 3 locations with the largest potential discrepancy reward peaks, L_D , and compared these predicted locations with expert-selected sampling location when they adapt their initial strategy. Fig. 8 illustrates two representative examples of the computed distribution of potential discrepancy, $D(l)$, for experts reporting different levels of confidence towards the given hypothesis. An expert, reporting high confidence towards the given hypothesis (Fig. 8A, C), selected location 16 (Fig. 8C, purple vertical line) to take measurements when adapt their initial sampling strategy. The selected location corresponds to a small potential discrepancy (Fig. 8C, green markers), which is consistent with the hypothesized reward, R_{validate} . In contrary, another expert, reporting low confidence towards the given hypothesis (Fig. 8B, D), selected location 10 (Fig. 8B, purple vertical line) to take measurements when adapt from their initial sampling strategy. The selected location corresponds to a large potential discrepancy (Fig. 8D, red markers), which is consistent with the hypothesized reward, $R_{\text{invalidate}}$.

Interestingly, the discrepancy based reward (Fig. 9, red markers) captured experts’ sampling location choices significantly better at high information level, *i.e.*, $I_{\text{total}} > 0.5$ (an average of 0.92 ± 0.07 accuracy), yet less effective at low information level (an average of 0.77 ± 0.21 accuracy). This is the opposite to the information based reward function (Fig. 9, blue markers), which was significantly more effective in capturing experts’ sampling behaviors at low information level, *i.e.*, $I_{\text{total}} < 0.5$ (an average of 0.91 ± 0.20 accuracy) than at high information level (an average of 0.59 ± 0.32 accuracy).

The clear crossover of high prediction rate, from information reward function at lower information level (Fig. 9, blue), to the discrepancy reward function at high information level (Fig. 9, red), suggested a transition of experts’ exploration focus, from an information coverage oriented objective to a hypothesis verification oriented objective, as the level of information increased. By simply applying the information based reward function for low information level (*i.e.*, $I_{\text{total}} < 0.5$), and discrepancy based reward function at high information level (*i.e.*, $I_{\text{total}} > 0.5$) could allow the robot to successfully capture the sampling behavior for the majority of experts.

6 CONCLUSION AND FUTURE DIRECTIONS

Our findings reveal that human scientists dynamically update science objectives in response to incoming measurements: when information is low, human scientists’ sampling decisions are mainly driven by an exploration objective, where they select sampling locations to efficiently increase spatial information coverage. Once scientists have collected sufficient amount of information to form an initial belief towards the hypothesis, they transition to a exploitation objective, where they select sampling locations to verify their beliefs about the hypothesis. The identified key decision factors allows for the extraction of two quantitative “reward functions”: an information based reward, and a potential discrepancy based reward. These objective based reward functions enable our field robotic assistants to infer experts’ dynamic sampling priorities during field data collection. By using an extremely simple model that predicts sampling locations based on the information based reward function at

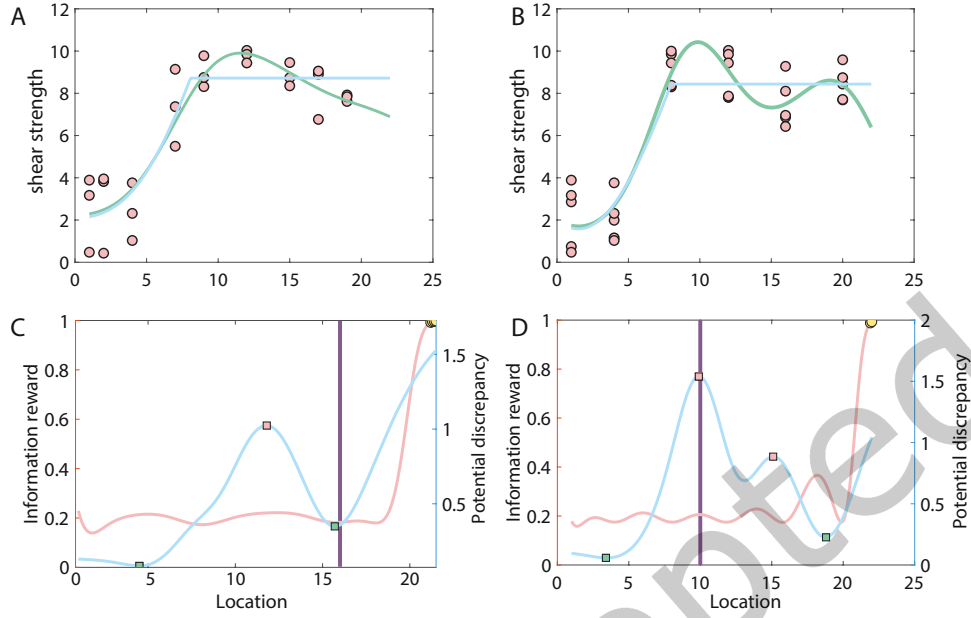


Fig. 8. An example illustration of reward distribution and sampling location choices for two experts. One expert, plotted in (A, C), reported high confidence towards the given hypothesis, while the other expert, plotted in (B, D), reported low confidence towards the given hypothesis. Red filled circles in (A, B) illustrate the existing measurements of shear strength that the experts have already sampled prior to the strategy deviation. Green curves represent the estimated shear strength, y , based on existing measurements. Blue curves represent the hypothesized shear strength, H_y , based on H_{A1} . (C, D) illustrate the spatial distribution of rewards, and expert-selected sampling location. Blue curve represents the potential discrepancy, with the red markers represent locations with large potential discrepancies, and the green markers represent sampling locations with small potential discrepancies. Red curve represents the information reward. Purple vertical line represents expert-selected location to sample next, upon deviation from their initial sampling strategy.

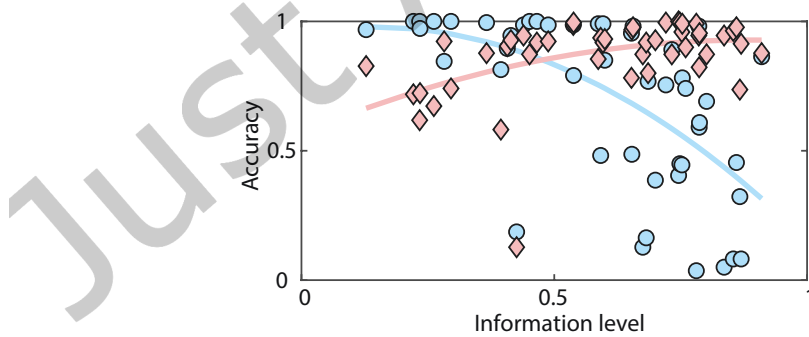


Fig. 9. Accuracy of robot-generated sampling locations at predicting expert geoscientist data collection decisions across different information coverage levels, based on the information reward (blue markers) and the discrepancy reward (red markers) functions. Blue and red curves are second order polynomial regression of markers with corresponding colors to provide visual aids of data trend.

low information level and the discrepancy based reward functions at high information level, the robot-predicted locations can successfully capture the sampling location choices for the majority of the human scientists.

The strikingly high prediction rate of the simple model reflects the success of our approach informed by the study of expert human behavior. Although we demonstrated our application with a geoscience problem, it is speculated that the identified decision factors and reward functions could be expanded to the study of other field sciences as well. Of course, human minds are extremely complicated, and there are many additional interesting questions that should be explored in order to achieve robust, disciplinary-agnostic robot-aided sampling. One interesting direction that future studies should further explore is decision making under more than one objective. In this study we noticed that at the information level around the transition point (*i.e.*, $I_{total} = 0.5$), the prediction accuracy with both information reward and discrepancy reward exhibited a temporary decrease. We suspect that this is because experts were balancing multiple objectives at this information range, and their exhibited sampling behaviors were reflective of different patterns of how humans balance multiple objectives. Future work should explore how experts select sampling locations to resolve multiple, and sometimes even conflicting, objectives. Another interesting direction is to expand the findings to more complex scenarios, such as natural environments with multiple spatial gradients, or scientific sampling under multiple competing alternative hypotheses. Application of the reward functions to these scenarios will require extracting hypotheses from scientists about how data is expected to change along the transect. Determining how data collection objectives change (or do not) with these different initial hypothesis forms is an important future research direction.

The understandings of human's science objectives and the connection to their sampling behaviors open up a number of avenues for future work on human-robot collaborative sampling. One possibility is to integrate our discoveries with adaptive planning and learning algorithms, to allow robots to suggest hypothesis-based sampling location adaptations to help lower scientists' cognitive load in low-level planing, freeing experts' minds to focus on scientific reasoning. Alternatively, these explicit decisions and reasoning patterns could help robots identify potential vulnerability to biases in humans' sampling decisions, and produce "nudges" to help correct human biases and enhance scientific outcomes. For example, we found that most experts had a propensity to chose sampling locations to confirm their beliefs rather than locations that were likely to contradict their beliefs. This confirmation bias could lead scientists to miss out on disconfirming evidence, leading to the pursuit of errant research directions and a waste of resources in those pursuits. A robot that can recognize confirmation bias – *i.e.*, recognize when scientists are collecting data with the objective of minimizing hypothesis discrepancy, despite poor hypothesis fit – could nudge scientists to reconsider their behavior by providing an explanation of why the inferred objective and corresponding behavior is suboptimal.

Another promising avenue is the development of cognitive-compatible robotic teammates that can effectively communicate with human teammates the tradeoffs in balancing multiple objectives, and collaboratively produce sampling decisions. The reward functions, extracted from human decision data set, could serve as a basic human decision models to allow robots to ask questions during the field sampling to help understand how human scientists balance multiple objectives. Over time, these understandings and human models will bring our robots beyond sensor instrumentation and closer to intelligent, trustworthy teammates, to help human scientists explore a wide variety of complex earth and planetary environments.

ACKNOWLEDGMENT

This work is supported by the NASA Planetary Science and Technology Through Analog Research (PSTAR) program, Award # 80NSSC22K1313. The authors would like to thank Shenyue Chen and Skyler Rankin for helping with the development of the web-based simulated scenario, Thomas F. Shipley and Daniel E. Koditschek for their support with human data collection, and Douglas J. Jerolmack, Ryan C. Ewing, Kenton R. Fisher, Marion Nachan for helpful discussions.

REFERENCES

- [1] Saurabh Arora and Prashant Doshi. 2021. A survey of inverse reinforcement learning: Challenges, methods and progress. *Artificial Intelligence* 297 (2021), 103500.
- [2] Todd S Braver. 2012. The variable nature of cognitive control: a dual mechanisms framework. *Trends in cognitive sciences* 16, 2 (2012), 106–113.
- [3] Alberto Candela, David Thompson, Eldar Noe Dobrea, and David Wettergreen. 2017. Planetary robotic exploration driven by science hypotheses for geologic mapping. In *2017 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 3811–3818.
- [4] Gregory Dudek, Philippe Giguere, Chris Prahacs, Shane Saunderson, Junaed Sattar, Luz-Abril Torres-Mendez, Michael Jenkin, Andrew German, Andrew Hogue, Arlene Ripsman, et al. 2007. Aqua: An amphibious autonomous robot. *Computer* 40, 1 (2007), 46–53.
- [5] Jonathan St BT Evans and Keith E Stanovich. 2013. Dual-process theories of higher cognition: Advancing the debate. *Perspectives on psychological science* 8, 3 (2013), 223–241.
- [6] Joseph L Fleiss, Bruce Levin, and Myunghee Cho Paik. 2013. *Statistical methods for rates and proportions*. John Wiley & sons.
- [7] Laura M Hiatt, Anthony M Harrison, and J Gregory Trafton. 2011. Accommodating human variability in human-robot teams through theory of mind. In *Twenty-Second International Joint Conference on Artificial Intelligence*.
- [8] Laura M Hiatt, Cody Narber, Esube Bekele, Sangeet S Khemlani, and J Gregory Trafton. 2017. Human modeling for human-robot collaboration. *The International Journal of Robotics Research* 36, 5-7 (2017), 580–596.
- [9] Trong Nghia Hoang, Bryan Kian Hsiang Low, Patrick Jaillet, and Mohan Kankanhalli. 2014. Nonmyopic ϵ -Bayes-optimal active learning of Gaussian processes. In *International Conference on Machine Learning*. PMLR, 739–747.
- [10] Geoffrey A Hollinger and Gaurav S Sukhatme. 2014. Sampling-based robotic information gathering algorithms. *The International Journal of Robotics Research* 33, 9 (2014), 1271–1287.
- [11] Jimin Hwang, Neil Bose, and Shuangshuang Fan. 2019. AUV adaptive sampling methods: A review. *Applied Sciences* 9, 15 (2019), 3145.
- [12] Hiroaki Inotsume, Takashi Kubota, and David Wettergreen. 2020. Robust Path Planning for Slope Traversing Under Uncertainty in Slip Prediction. *IEEE Robotics and Automation Letters* 5, 2 (2020), 3390–3397. <https://doi.org/10.1109/LRA.2020.2975756>
- [13] Oussama Khatib, Xiyang Yeh, Gerald Brantner, Brian Soe, Boyeon Kim, Shameek Ganguly, Hannah Stuart, Shiquan Wang, Mark Cutkosky, Aaron Edsinger, et al. 2016. Ocean one: A robotic avatar for oceanic discovery. *IEEE Robotics & Automation Magazine* 23, 4 (2016), 20–29.
- [14] Andreas Krause, Ajit Singh, and Carlos Guestrin. 2008. Near-optimal sensor placements in Gaussian processes: Theory, efficient algorithms and empirical studies. *Journal of Machine Learning Research* 9, 2 (2008).
- [15] Siyuan Liu, Miguel Araujo, Emma Brunskill, Rosaldo Rossetti, Joao Barros, and Ramayya Krishnan. 2013. Understanding sequential decisions via inverse reinforcement learning. In *2013 IEEE 14th International Conference on Mobile Data Management*, Vol. 1. IEEE, 177–186.
- [16] Yong Ma, Mengqi Hu, and Xinpeng Yan. 2018. Multi-objective path planning for unmanned surface vehicle with currents effects. *ISA Transactions* 75 (2018), 137–156. <https://doi.org/10.1016/j.isatra.2018.02.003>
- [17] Roman Marchant and Fabio Ramos. 2014. Bayesian Optimisation for informative continuous path planning. In *2014 IEEE International Conference on Robotics and Automation (ICRA)*. 6136–6143. <https://doi.org/10.1109/ICRA.2014.6907763>
- [18] C McKenna-Neuman and WG Nickling. 1989. A theoretical and wind tunnel investigation of the effect of capillary water on the entrainment of sediment by wind. *Canadian Journal of Soil Science* 69, 1 (1989), 79–96.
- [19] Tim Miller. 2019. Explanation in artificial intelligence: Insights from the social sciences. *Artificial Intelligence* 267 (2019), 1–38.
- [20] Wim De Neys. 2006. Dual processing in reasoning: Two systems but one reasoner. *Psychological science* 17, 5 (2006), 428–433.
- [21] Petter Nilsson, Sofie Haesaert, Rohan Thakker, Kyohei Otsu, Cristian-Ioan Vasile, Ali-Akbar Agha-Mohammadi, Richard M Murray, and Aaron D Ames. 2018. Toward specification-guided active mars exploration for cooperative robot teams. (2018).
- [22] EZ Noe Dobrea, C Ahrens, ME Banks, A Breitfeld, A Candela, RN Clark, M Hansen, A Hendrix, G Holsclaw, G Kramer, et al. 2022. Autonomous Rover Science in the Field: First Results. *LPI Contributions* 2678 (2022), 1674.
- [23] Arvind Pereira, Jonathan Binney, Geoffrey A. Hollinger, and Gaurav S. Sukhatme. 2013. Risk-aware Path Planning for Autonomous Underwater Vehicles using Predictive Ocean Models. *Journal of Field Robotics* 30 (2013).
- [24] G. Picardi, M. Chellapurath, S. Iacoponi, S. Stefanni, C. Laschi, and M. Calisti. 2020. Bioinspired underwater legged robot for seabed exploration with low environmental disturbance. *Science Robotics* 5, 42 (2020), eaaz1012. <https://doi.org/10.1126/scirobotics.aaz1012> arXiv:<https://www.science.org/doi/pdf/10.1126/scirobotics.aaz1012>
- [25] Feifei Qian, Douglas Jerolmack, Nicholas Lancaster, George Nikolich, Paul Reverdy, Sonia Roberts, Thomas Shipley, R Scott Van Pelt, Ted M Zobeck, and Daniel E Koditschek. 2017. Ground robotic measurement of aeolian processes. *Aeolian research* 27 (2017), 1–11.
- [26] Feifei Qian, Dylan Lee, George Nikolich, Daniel Koditschek, and Douglas Jerolmack. 2019. Rapid In Situ Characterization of Soil Erodibility With a Field Deployable Robot. *Journal of Geophysical Research: Earth Surface* 124, 5 (2019), 1261–1280.
- [27] Alberto Quattrini Li. 2020. Exploration and mapping with groups of robots: Recent trends. *Current Robotics Reports* 1, 4 (2020), 227–237.
- [28] Uluc Saranlı, Martin Buehler, and Daniel E Koditschek. 2001. RHex: A simple and highly mobile hexapod robot. *The International Journal of Robotics Research* 20, 7 (2001), 616–631.

- [29] Matthias Seeger. 2004. Gaussian processes for machine learning. *International journal of neural systems* 14, 02 (2004), 69–106.
- [30] Florian Shkurti, Anqi Xu, Malika Meghjani, Juan Camilo, Juan Higuera, Yogesh Girdhar, Philippe Giguère, Bir Dey, Jimmy Li, Arnold Kalmbach, Chris Prahacs, Katrine Turgeon, Ioannis Rekleitis, and Gregory Dudek. 2012. Multi-Domain Monitoring of Marine Environments using a Heterogeneous Robot Team. <https://doi.org/10.1109/IROS.2012.6385685>
- [31] Thane Somers and Geoffrey A Hollinger. 2016. Human–robot planning and learning for marine data collection. *Autonomous Robots* 40, 7 (2016), 1123–1137.
- [32] Sebastian Thrun et al. 2002. Robotic mapping: A survey. *Exploring artificial intelligence in the new millennium* 1, 1-35 (2002), 1.
- [33] Jonatan Scharff Willners, Lachlan Toohey, and Yvan Petillot. 2019. Sampling-Based Path Planning for Cooperative Autonomous Maritime Vehicles to Reduce Uncertainty in Range-Only Localization. *IEEE Robotics and Automation Letters* 4, 4 (2019), 3987–3994. <https://doi.org/10.1109/LRA.2019.2926947>
- [34] Cristina G Wilson, Clare Elizabeth Bond, and Thomas F Shipley. 2019. How can geologic decision making under uncertainty be improved? *Solid earth* (2019), 1–34.
- [35] Cristina G Wilson, Feifei Qian, Douglas J Jerolmack, Sonia Roberts, Jonathan Ham, Daniel Koditschek, and Thomas F Shipley. 2021. Spatially and temporally distributed data foraging decisions in disciplinary field science. *Cognitive Research: Principles and Implications* 6, 1 (2021), 1–16.
- [36] Ciyu Zhu, Richard H. Byrd, Peihuang Lu, and Jorge Nocedal. 1997. Algorithm 778: L-BFGS-B: Fortran Subroutines for Large-Scale Bound-Constrained Optimization. *ACM Trans. Math. Softw.* 23, 4 (dec 1997), 550–560. <https://doi.org/10.1145/279232.279236>