

Adaptive Sampling for Seabed Identification from Ambient Acoustic Noise

Matthew Sullivan

*Dept. of Electrical and Computer Engineering
Portland State University
matsu2@pdx.edu*

John Gebbie

*Metron, Inc.
gebbie@metsci.com*

John Lipor

*Dept. of Electrical and Computer Engineering
Portland State University
lipor@pdx.edu*

Abstract—We study the problem of adaptively obtaining ambient acoustic measurements via an autonomous underwater vehicle, with the goal of characterizing the geoacoustic properties of the seabed. In contrast to the traditional adaptive sampling scenario, we are provided with sets of snapshots associated with each spatial location, making the problem one of unsupervised learning. We demonstrate how sets of snapshots can be used to obtain noisy pairwise similarities between locations, which can then be used to perform level set estimation to separate the seabed into two highly-distinct types. We propose an adaptive sampling policy that aims to directly reduce the number of locations whose level set membership is uncertain, as well as an approach to minimizing the distance traveled while sampling. Results on synthetic and real-world sediment data demonstrate the benefits of our approach in terms of both accuracy and distance traveled.

I. INTRODUCTION

Understanding the spatial variability of geoacoustic properties in the ocean is essential for characterizing SONAR system performance [1]–[3]. Historically, spatial variability has been estimated through collection of sediment cores or active SONAR, but these costly procedures are only feasible for relatively small regions of the ocean. Further, existing approaches collect measurements opportunistically or on a uniform grid, rather than letting previous measurements guide the sampling procedure to the most informative locations. Recently, researchers have demonstrated the potential to estimate seabed parameters through ambient acoustic noise sources, such as surface waves or passing ships [4]–[6]. Such sources have the potential to enable low-powered autonomous underwater vehicles (AUVs) to continuously capture information over vast regions of the ocean.

Given the massive scale of the regions of interest to practitioners, obtaining a high-fidelity estimate of the seabed requires adaptively guiding AUVs to discover and track regions where the seabed varies. In recent years, a wide variety of adaptive sampling algorithms have been developed to improve environmental monitoring, improving estimates of phenomena such as thermoclines [7]–[9], concentrations of harmful chemicals and bacteria [10], [11], and wildfire boundaries [12], [13]. Such algorithms typically proceed by defining a notion of information gain, which is balanced with the cost of obtaining measurements and traveling to new sampling locations.

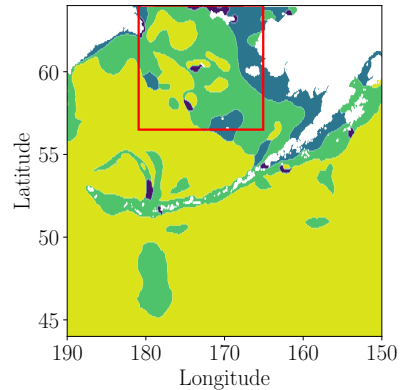


Fig. 1: Seabed types (by color) in the northern Pacific Ocean. The seabed predominantly consists of clay (yellow) or medium silt (green), and the goal is to discover all locations where the type differs significantly (light/dark blue). The red box denotes the subregion used for experiments in Sec. IV.

Existing adaptive sampling algorithms typically proceed by either a local view, where the goal is to track a single boundary as closely as possible [7]–[9], [12]–[16], or a global view, in which the entire field of interest is regressed after each measurement [10], [11], [17]–[23]. While the former are typically more efficient and noise-tolerant, they are limited to the case of a single boundary and have no principled means of exploring the region of interest. In the case of seabed characterization, the goal is to discover multiple disconnected regions where the seabed differs from the dominant type (see Fig. 1). However, global approaches typically perform poorly in high-noise scenarios, with a tendency to over-explore. Further, in the case of seabed characterization, direct measurements of seabed type are not available. Instead, the AUV continuously collects acoustic snapshots, which can be associated to a spatial location. As a result, seabed characterization can be viewed as an unsupervised problem, where the goal is to simultaneously decide which snapshots correspond to the same seabed type and determine the most efficient path for the AUV.

In this work, we propose an approach to performing adaptive sampling based on the collection of ambient acoustic snapshots. We show how the seabed characterization problem can instead be formulated as one of level set estimation

(LSE) from noisy similarity measurements. We then propose an adaptive sampling policy that seeks to directly reduce the number of uncertain locations (i.e., locations whose level set membership is unknown). Finally, we introduce a mechanism for controlling the distance traveled by the vehicle that has the ancillary benefit of reducing the computation time required by our policy. We show that our approach outperforms existing adaptive sampling policies on both synthetic data and a case study of seabed characterization in the northern Pacific Ocean.

II. PROBLEM FORMULATION

Consider an AUV carrying an array of M receivers that capture ambient acoustic noise, typically assumed to be generated by surface waves. The time-series pressure recordings are Fourier transformed to obtain snapshots $z \in \mathbb{C}^M$ at a single frequency, which are assumed to be drawn from a circularly-symmetric complex Gaussian distribution with covariance [5]

$$\Sigma_\theta = \mathbb{E}[zz^H] = \sigma_s^2 \Gamma_\theta + \sigma_n^2 I, \quad (1)$$

where σ_s^2 is the power in the ambient noise, σ_n^2 is the non-acoustic sensor noise variance, and Γ_θ is the signal covariance matrix, with parameters θ corresponding to the physical properties of the seabed [24], [25]. For a spatial location $x \in \mathbb{R}^2$, there exists a corresponding seabed type defined by its covariance. For two locations x, x' , our goal is to determine whether the snapshots are generated from the same (or a very similar) covariance.

We first motivate our LSE formulation with an application of interest. The Naval Oceanographic Office maintains a bottom sediment type database with a list of 23 seabed types provided in the High Frequency Environmental Acoustics (HFEVA) dataset [26]. In practice, the seabed often consists of one predominant *background* type (e.g., clay), with patches of different types distributed throughout a region (see Fig. 1). Hence, the chief practical goal is to determine all regions where the seabed differs significantly from this background type. In the context of Fig. 1, the background types clay (yellow) and medium silt (green) are extremely similar, making the goal to discover regions corresponding to rock (dark blue) and gravelly muddy sand (light blue).

To formulate the above problem in terms of LSE, we consider a finite domain of interest $\mathcal{D} \subset \mathbb{R}^2$ and select a reference location x_0 known to belong to the background sediment type. Observe that sampling location x_t at time t provides a collection of L snapshots $z_1^{(t)}, \dots, z_L^{(t)}$. These snapshots can be used to obtain similarities between pairs of locations as follows. For each sampled location x_t , we compute the sample covariance matrix $\hat{\Sigma}_t$ of the corresponding snapshots. We then compute the Jensen-Shannon divergence (JSD) [27] between locations x_t and x_0

$$J(\hat{\Sigma}_0 || \hat{\Sigma}_t) = (D(\hat{\Sigma}_0 || \hat{\Sigma}_t) + D(\hat{\Sigma}_t || \hat{\Sigma}_0))/2, \quad (2)$$

where $D(\hat{\Sigma}_0 || \hat{\Sigma}_t)$ is the Kullback-Leibler divergence between the Gaussian distributions defined by the covariance estimates

$\hat{\Sigma}_0$ and $\hat{\Sigma}_t$ [28]. We exponentiate the JSD to obtain the noisy similarity between x_0 and x_t

$$s_t = \exp\left(-J(\hat{\Sigma}_0 || \hat{\Sigma}_t)/\ell^2\right), \quad (3)$$

where ℓ is a tuning parameter used to control the scale of the similarities. Our goal is to accurately estimate the sublevel set of locations that are sufficiently dissimilar to x_0

$$\mathcal{L} = \{x \in \mathcal{D} : s(x_0, x) \leq \tau\}, \quad (4)$$

where $\tau > 0$ is the threshold governing the degree of dissimilarity of interest and $s(x_0, x)$ is the similarity between the locations using the unknown true covariance matrices. In addition, we wish to minimize the total sampling time, which depends primarily on the distance traveled.

III. PROPOSED SAMPLING POLICY

In this section, we describe our approach to distance-penalized LSE. As with other global algorithms, our policy sequentially updates a prediction and confidence interval at each $x \in \mathcal{D}$, as well as estimates of the sublevel set, superlevel set, and the uncertain set, which consists of points whose level set membership cannot be confidently determined. Our key observation is that the goal of LSE is to reduce the cardinality of the uncertain set as quickly as possible. In light of this, our proposed sampling policy chooses the location that obtains the greatest (approximate) reduction in the size of the uncertain set. Further, a great deal of attention has been given to various forms of distance penalization in adaptive sampling. We account for this cost by limiting our policy to only select from among a fixed number of nearest neighbors within the uncertain set. This directly reduces the distance traveled, while selecting only from the uncertain set ensures that the selected point still provides significant information gain.

To obtain predictions and confidence intervals, we utilize kernel ridge regression (KRR), though we note that our approach works with any regression algorithm that admits an uncertainty estimate. Let $k(x, x')$ be a kernel function, which denotes the similarity between locations x and x' . Consider a set of visited sample locations $x_1, \dots, x_t$ with corresponding similarity measurements $s_1, \dots, s_t$. Define $K_t \in \mathbb{R}^{t \times t}$ to be the matrix whose i, j th entry is $k(x_i, x_j)$, let $k_{x,t} = [k(x, x_1), \dots, k(x, x_t)]^T$, and $y_t = [s_1, \dots, s_t]^T$. For a given point $x \in \mathcal{D}$, the KRR prediction and confidence width are

$$\mu_t(x) = k_{x,t}^T (K_t + \gamma I)^{-1} y_t \quad (5)$$

and

$$\sigma_t(x) = \gamma^{-1/2} \sqrt{k(x, x) - k_{x,t}^T (K_t + \gamma I)^{-1} k_{x,t}}, \quad (6)$$

where $\gamma \geq 0$ is a regularization parameter. Following [29], [30], these predictions can be updated in an online fashion to reduce computational complexity.

Based on the above predictions and confidence widths, we can state with high probability that the true similarity lies within the confidence interval $C_t(x) = [\mu_t(x) \pm \eta \sigma_t(x)]$, where $\eta > 0$ controls the size of the confidence interval. While

Algorithm 1 Lookahead Uncertain Set Reduction (LUSR)

Input: domain $\mathcal{D}$, initial location x_0 , kernel function k , threshold τ , tuning parameters γ, η , number of neighbors ρ , number of samples T

Output: predicted sets L_T, H_T

- 1: $L_0 \leftarrow \emptyset, H_0 \leftarrow \emptyset, U_0 \leftarrow \mathcal{D}$
- 2: **for** $t = 0, \dots, T$ **do**
- 3: $\mathcal{N}_\rho(x_t) \leftarrow \rho$ nearest neighbors of x_t in U_t
- 4: $x_{t+1} \leftarrow \arg \max_{x \in \mathcal{N}_\rho(x_t)} |U_t| - |\tilde{U}_{t+1}|$
- 5: collect snapshots, form similarity s_{t+1} according to (3)
- 6: update μ_{t+1} and σ_{t+1} according to (5)-(6)
- 7: update $L_{t+1}, H_{t+1}, U_{t+1}$ according to (7)-(9)
- 8: **end for**

the value of η can be motivated theoretically, in practice most approaches select a constant value that works well over a range of problems [17]. Based on the confidence interval above, we maintain estimates of the sublevel, superlevel, and uncertain sets, respectively

$$L_t = \{x \in \mathcal{D} : \mu_t(x) + \eta\sigma_t(x) \leq \tau\} \quad (7)$$

$$H_t = \{x \in \mathcal{D} : \mu_t(x) - \eta\sigma_t(x) \geq \tau\} \quad (8)$$

$$U_t = \mathcal{D} \setminus (L_t \cup H_t). \quad (9)$$

We now define our sampling policy, which we refer to as *lookahead uncertainty set reduction* (LUSR). As stated above, the goal of LUSR is to select the sample that results in the greatest reduction of the uncertain set. Note that the confidence width (6) does not depend on the value of the measurement and can therefore be evaluated exactly prior to sampling a given point, as utilized by [11], [31]. However, determining the uncertain set also requires evaluating the posterior prediction (5), which does depend on the measurement value. While this value cannot be known without sampling, our key observation is that we can utilize the existing confidence interval to obtain a false measurement for a given location. In particular, we set

$$\tilde{s}_{t+1} = \begin{cases} \mu_t(x) - \eta\sigma_t(x), & \mu_t(x) > \tau \\ \mu_t(x) + \eta\sigma_t(x), & \mu_t(x) \leq \tau \end{cases} \quad (10)$$

and then form estimates $\tilde{L}_{t+1}$, $\tilde{H}_{t+1}$, and $\tilde{U}_{t+1}$ based on this value. LUSR then selects the sample that maximizes the estimated uncertain set reduction

$$x_{t+1} = \arg \max_{x \in U_t} |U_t| - |\tilde{U}_{t+1}|. \quad (11)$$

Pseudocode for sampling with LUSR is given in Alg. 1.

The false measurement (10) can be viewed as a pessimistic choice in the sense that it is the value that removes the fewest points from the uncertain set. If $\mu_t(x) > \tau$, we set $\tilde{s}_{t+1}$ to be the smallest possible value, which pushes x and all similar points closer to the level set boundary τ . While one could utilize any value within $C_t(x)$, our initial experimentation showed that pessimism results in a lower distance traveled than other options, since it guides the vehicle toward locations where we are highly confident many points will be removed.

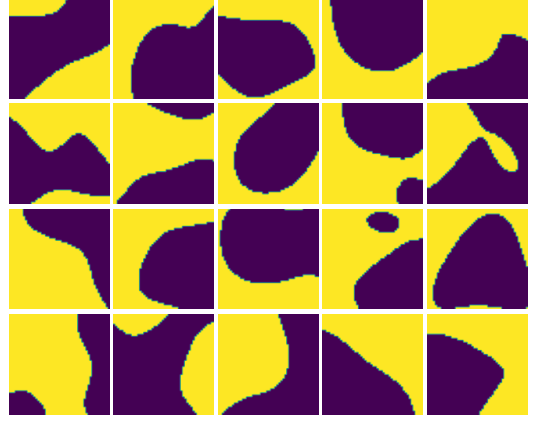


Fig. 2: Synthetic fields used for evaluation of LSE policies.

Finally, we describe our approach to distance penalization, which also reduces the computational complexity of our approach. To reduce the distance traveled, we wish to avoid large jumps across the region of interest. One way to prevent such jumps is to only allow the policy to select from among the ρ nearest neighbors within the uncertain set. At time t , we denote this set by $\mathcal{N}_\rho(x_t)$. We then perform the maximization in (11) over $\mathcal{N}_\rho(x_t)$, rather than over the entirety of U_t . We show empirically in Sec. IV that there is a tradeoff between distance traveled and accuracy that is governed by ρ . As an additional benefit, searching over a fixed number of neighbors can dramatically reduce the computation cost of LUSR, especially for the case of large domains. Thus, our approach provides a simple means to distance penalization that also makes lookahead-type methods such as ours feasible for real-time adaptive sampling.

IV. EXPERIMENTAL RESULTS

In this section, we evaluate the empirical performance of LUSR on both synthetic data and real types maps from the HFEVA database. We compare our policy to random sampling, margin-based sampling $x_{mar} = \arg \min_x |\mu_t(x) - \tau|$, as well as max variance sampling (VAR) $x_{var} = \arg \max_x \sigma_t(x)$. For all policies, we perform distance penalization by searching over the ρ nearest neighbors. Although not shown, we also compared against the state-of-the-art straddle heuristic [32] and variance reduction in the spirit of [11]. However, the straddle heuristic performed similar to margin while requiring selection of another tuning parameter, while variance reduction was nearly equivalent to maximum variance sampling.

We first consider synthetic data, generating 20 random fields of size 50×50 by mapping the coordinates to a fourth-order polynomial, generating a random weight vector in this high-dimensional space, then labeling each location according to a linear classifier defined by this weight vector. The resulting fields are shown in Fig. 2, where yellow denotes the superlevel set and blue denotes the sublevel set. For each field, we perform 32 random trials, drawing similarities from a truncated normal distribution with means zero and one for within-class and across-class similarities, respectively.

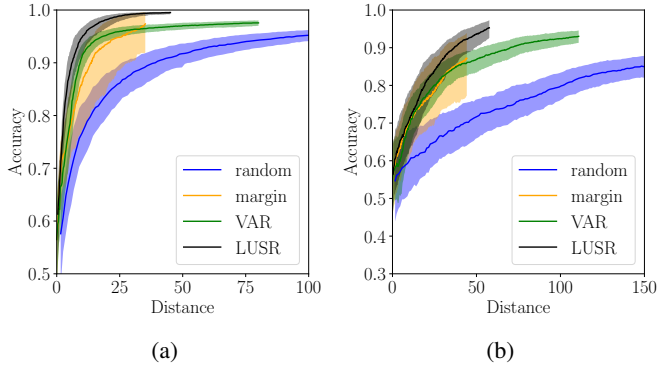


Fig. 3: Level set estimation accuracy versus distance traveled on synthetic data. (a) Low-noise setting with standard deviation 0.15. (b) High-noise setting with standard deviation 0.45.

We set the level-set threshold $\tau = 0.5$ and consider a low-noise scenario having standard deviation 0.15 and a high-noise scenario with standard deviation 0.45. For all policies, we set $\gamma = 0.01$, $\eta = 0.5$, the number of neighbors $\rho = 100$, and use the top left corner as the reference location. We evaluate algorithms based on the classification accuracy, treating the level set membership as a binary label.

Fig. 3 shows the median accuracy versus distance traveled for each policy considered in both the low-noise and high-noise settings, along with the interquartile range. In both cases, we see that LUSR achieves a higher accuracy than its closest competitor at a lower distance traveled. We note that for fields containing only a single boundary, margin is able to quickly discover and track the boundary, achieving a high accuracy at a small cost. However, for fields whose super/sublevel sets consist of more than one component, margin lacks the exploration benefits of LUSR and obtains a low accuracy. While VAR quickly discovers all components, it focuses too heavily on exploration, resulting in a much higher distance traveled. Hence, we see that LUSR selects points that are both near the level set boundary and of high uncertainty, obtaining the appropriate degree of exploration.

Next, we evaluate our approach to distance penalization. Fig. 4 shows the accuracy and distance traveled as a function of the number of neighbors for the low-noise setting described above. As expected, by limiting the number of neighbors, we can obtain a significant reduction in distance traveled by all sampling policies at the cost of estimation accuracy. Further, we again see that LUSR obtains a balance between the high accuracy of VAR while traveling distance on par with margin.

Finally, we compare algorithm performance on realistic sediment type data from the region bounded by the red box in Fig. 1. This region contains HFEVA sediment types 2, 8, 16, 17, and 22, with types 16-22 having very similar covariances [33]. Hence, we aim to distinguish types 2 and 8 from the background types 16, 17, and 22. For each sediment type, we obtain a signal covariance from the Multidimensional Ambient Noise Model [34] and generate complex normal snapshots according to the covariance defined by (1) with a signal-to-noise ratio of 10, using $M = 32$ sensors. Each sampled

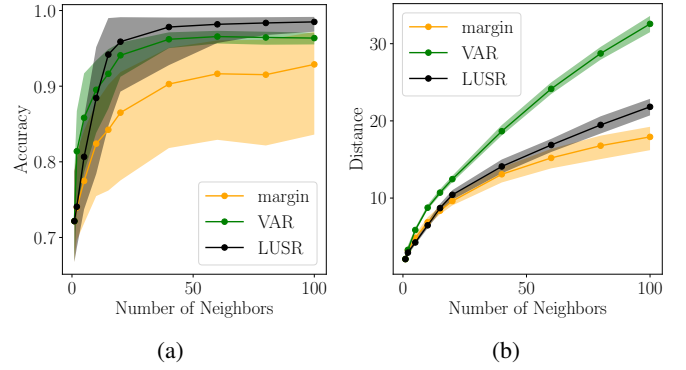


Fig. 4: Impact of number of neighbors considered on (a) accuracy and (b) distance traveled for different selection policies. The figures show a trade-off between accuracy and distance traveled that is compatible with any policy.

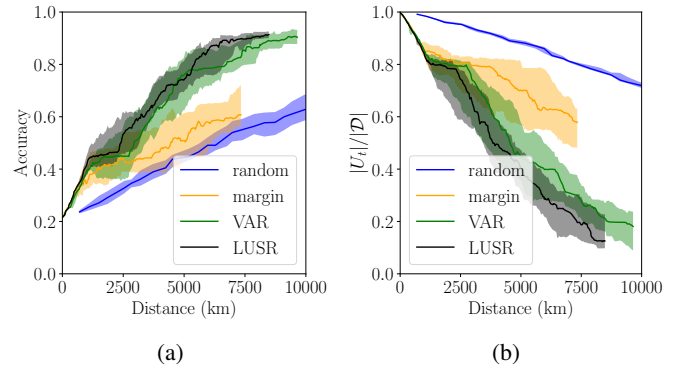


Fig. 5: Level set estimation accuracy on seabed data from the northern Pacific Ocean. (a) Accuracy as a function of distance. (b) Relative size of uncertain set as a function of distance.

location provides $L = 300$ snapshots. To obtain similarities, we set $\ell = 10$, and the level set threshold is chosen to be $\tau = 0.3$, though all policies are robust to a wide range of values for these parameters. Results are shown for 16 random trials with the bottom left corner as the reference location.

Fig. 5 displays (a) the accuracy as well as (b) the relative size of the uncertain set as a function of distance. As with synthetic data, we see that LUSR achieves a higher accuracy than VAR at a lower cost, with a maximum distance approximately 1,200 km less than that required by VAR. Further, Fig. 5(b) shows that LUSR reduces the uncertain set more quickly than VAR, indicating that our use of false measurements obtains the desired benefit of estimating the cardinality of U_{t+1} .

V. CONCLUSION

We have provided an approach to seabed characterization from ambient acoustic noise through the lens of adaptive sampling for level set estimation. Our proposed adaptive sampling policy directly estimates the reduction in uncertainty via the use of false measurements, obtaining strong empirical results on both synthetic and realistic ambient acoustic data. The analysis of both our adaptive sampling policy and our distance penalization technique is of particular interest as a topic of future study.

REFERENCES

- [1] E. L. Hamilton, "Geoacoustic models of the sea floor," *Physics of sound in marine sediments*, pp. 181–221, 1974.
- [2] D. DEL BALZO, "Critical angle and seabed scattering issues for active-sonar performance predictions in shallow water," in *High Frequency Acoustics in Shallow Water, Conference Proceedings, 1997*, 1997.
- [3] M. Prior, C. Harrison, and S. Healy, "Assessment of the impact of uncertainty in seabed geoacoustic parameters on predicted sonar performance," *Impact of Littoral Environmental Variability of Acoustic Predictions and Sonar Performance*, pp. 531–538, 2002.
- [4] L. Muzi, M. Siderius, and C. M. Verlinden, "Passive bottom reflection-loss estimation using ship noise and a vertical line array," *The Journal of the Acoustical Society of America*, vol. 141, no. 6, pp. 4372–4379, 2017.
- [5] J. Gebbie and M. Siderius, "Optimal environmental estimation with ocean ambient noise," *The Journal of the Acoustical Society of America*, vol. 149, no. 2, pp. 825–834, 2021.
- [6] M. Siderius and J. Gebbie, "Environmental information content of ocean ambient noise," *The Journal of the Acoustical Society of America*, vol. 146, no. 3, pp. 1824–1833, 2019.
- [7] S. Petillo, H. Schmidt, P. Lermusiaux, D. Yoerger, and A. Balasuriya, "Autonomous & adaptive oceanographic front tracking on board autonomous underwater vehicles," in *OCEANS 2015-Genova*. IEEE, 2015, pp. 1–10.
- [8] J. Lipor and G. Dasarthy, "Quantile search with time-varying search parameter," in *2018 52nd Asilomar Conference on Signals, Systems, and Computers*. IEEE, 2018, pp. 1016–1018.
- [9] D. Wang, G. Dasarthy, and J. Lipor, "Distance-penalized active learning via markov decision processes," in *Proc. IEEE Data Science Workshop*, 2019.
- [10] G. Hitz, A. Gotovos, M.-É. Garneau, C. Pradalier, A. Krause, R. Y. Siegwart *et al.*, "Fully autonomous focused exploration for robotic environmental monitoring," in *Robotics and Automation (ICRA), 2014 IEEE International Conference on*. IEEE, 2014, pp. 2658–2664.
- [11] I. Bogunovic, J. Scarlett, A. Krause, and V. Cevher, "Truncated variance reduction: A unified approach to bayesian optimization and level-set estimation," in *Advances in neural information processing systems*, 2016, pp. 1507–1515.
- [12] P. Kearns, J. Lipor, and B. Jedynek, "Optimal adaptive sampling for boundary estimation with mobile sensors," in *2019 53rd Asilomar Conference on Signals, Systems, and Computers*. IEEE, 2019, pp. 1621–1625.
- [13] S. Stalley and J. Lipor, "A graph-based approach to boundary tracking with mobile sensors," *IEEE Robotics and Automation Letters*, 2021, accepted for publication. [Online]. Available: <http://web.cecs.pdx.edu/~lipor/Papers/stalley2022graph.pdf>
- [14] Z. Jin and A. L. Bertozzi, "Environmental boundary tracking and estimation using multiple autonomous vehicles," in *2007 46th IEEE conference on decision and control*. IEEE, 2007, pp. 4918–4923.
- [15] A. Joshi, T. Ashley, Y. R. Huang, and A. L. Bertozzi, "Experimental validation of cooperative environmental boundary tracking with on-board sensors," in *2009 American Control Conference*. IEEE, 2009, pp. 2630–2635.
- [16] G. Dasarthy, R. Nowak, and X. Zhu, "S2: An efficient graph based active learning algorithm with application to nonparametric classification," in *Conference on Learning Theory*, 2015, pp. 503–522.
- [17] A. Gotovos, N. Casati, G. Hitz, and A. Krause, "Active learning for level set estimation," in *IJCAI*, 2013, pp. 1344–1350.
- [18] L. Bottarelli, J. Blum, M. Bicego, and A. Farinelli, "Path efficient level set estimation for mobile sensors," in *Proceedings of the Symposium on Applied Computing*, 2017, pp. 262–267.
- [19] D. LeJeune, G. Dasarthy, and R. Baraniuk, "Thresholding graph bandits with graphl," in *International Conference on Artificial Intelligence and Statistics*. PMLR, 2020, pp. 2476–2485.
- [20] P. Thaker, M. Malu, N. Rao, and G. Dasarthy, "Maximizing and satisficing in multi-armed bandits with graph information," *Advances in Neural Information Processing Systems*, vol. 35, pp. 2019–2032, 2022.
- [21] B. Mason, L. Jain, S. Mukherjee, R. Camilleri, K. Jamieson, and R. Nowak, "Nearly optimal algorithms for level set estimation," in *International Conference on Artificial Intelligence and Statistics*. PMLR, 2022, pp. 7625–7658.
- [22] J. P. Folch, S. Zhang, R. M. Lee, B. Shafei, D. Walz, C. Tsay, M. van der Wilk, and R. Misener, "Snake: Bayesian optimization via pathwise exploration," in *2022 AIChE Annual Meeting*. AIChE, 2022.
- [23] S. S. Ramesh, P. G. Sessa, A. Krause, and I. Bogunovic, "Movement penalized bayesian optimization with application to wind energy systems," *Advances in Neural Information Processing Systems*, vol. 35, pp. 27 036–27 048, 2022.
- [24] W. A. Kuperman and F. Ingenito, "Spatial correlation of surface generated noise in a stratified ocean," *The journal of the acoustical society of America*, vol. 67, no. 6, pp. 1988–1996, 1980.
- [25] C. Harrison, "Formulas for ambient noise level and coherence," *The journal of the acoustical society of America*, vol. 99, no. 4, pp. 2055–2066, 1996.
- [26] Naval Oceanographic Office Acoustics Division, "Database description for bottom sediment type (U)," 2003.
- [27] J. Lin, "Divergence measures based on the shannon entropy," *IEEE Transactions on Information theory*, vol. 37, no. 1, pp. 145–151, 1991.
- [28] S. Kullback and R. A. Leibler, "On information and sufficiency," *The annals of mathematical statistics*, vol. 22, no. 1, pp. 79–86, 1951.
- [29] F. Zhang, *The Schur complement and its applications*. Springer Science & Business Media, 2006, vol. 4.
- [30] M. Valko, N. Korda, R. Munos, I. Flounas, and N. Cristianini, "Finite-time analysis of kernelised contextual bandits," in *Uncertainty in Artificial Intelligence*, 2013.
- [31] P. Hennig and C. J. Schuler, "Entropy search for information-efficient global optimization," *Journal of Machine Learning Research*, vol. 13, no. 6, 2012.
- [32] B. Bryan, R. C. Nichol, C. R. Genovese, J. Schneider, C. J. Miller, and L. Wasserman, "Active learning for identifying function threshold boundaries," *Advances in neural information processing systems*, vol. 18, 2005.
- [33] J. Lipor, J. Gebbie, and M. Siderius, "On the limits of distinguishing seabed types via ambient acoustic sound," *The Journal of the Acoustical Society of America*, 2023, accepted for publication.
- [34] Naval Oceanographic Office Acoustics Division, "Software design description and software test description for the multi-dimensional ambient noise model (MDANM) version 1.01 (U)," 2022.