

In Situ Analysis for Intelligent Control

Maria Fox and Derek Long
Dept. Computer and Information Sciences,
University of Strathclyde, Glasgow, UK
Email: {maria,derek}@cis.strath.ac.uk

Frederic Py and Kanna Rajan and John Ryan
Monterey Bay Aquarium Research Institute
Moss Landing, USA
Email: {fpy,kanna,ryjo}@mbari.org

Abstract—We report a pilot study on *in situ* analysis of backscatter data for intelligent control of a scientific instrument on an Autonomous Underwater Vehicle (AUV) carried out at the Monterey Bay Aquarium Research Institute (MBARI). The objective of the study is to investigate techniques which use machine intelligence to enable event-response scenarios. Specifically we analyse a set of techniques for automated sample acquisition in the water-column using an electro-mechanical “Gulper”, designed at MBARI. This is a syringe-like sampling device, carried onboard an AUV. The techniques we use in this study are clustering algorithms, intended to identify the important distinguishing characteristics of bodies of points within a data sample. We demonstrate that the complementary features of two clustering approaches can offer robust identification of interesting features in the water-column, which, in turn, can support automatic event-response control in the use of the Gulper.

I. INTRODUCTION

Understanding processes in the complex coastal ocean is very challenging due to the effects of diverse forces from the atmosphere, ocean circulation, and land-sea interactions. The coastal environment and its resources are influenced by fluctuations extending over a vast range of time and space scales, from global-scale multidecadal variability [2] to small-scale episodic events [17]. Under-sampling remains a primary limitation in coastal oceanography, and this limitation is being addressed by recent advancements in autonomous underwater vehicles (AUVs) that permit frequent, high-resolution mapping of coastal waters. Novel sensors and diverse sensor suites on AUVs are advancing the spectrum of measurements needed to solve complex, interdisciplinary problems. While sensor advancements are a great boon to research, some research requires analysis of water samples in laboratories on shore. To cope with these concerns, Monterey Bay Aquarium Research Institute (MBARI) has designed the Gulper to be embedded on an AUV and, in principle, given the ability to autonomously take water samples based on *in-situ* sensor inputs. The vast data archive from this AUV, now spanning four years of frequent surveys, clearly emphasizes the extremely episodic nature of variability in the region. Processes of great scientific interest occur regularly, but the temporal and spatial attributes of many processes are highly unpredictable. The intersection of limited predictability in processes and a highly capable AUV motivates the development of machine learning capabilities that permit the AUV to recognize and sample specific environmental features.

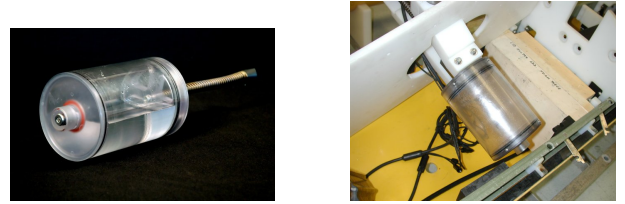


Fig. 1. The Gulper water-sampling device and its mounting inside an AUV mid-section.

We report a pilot study on *in situ* analysis of backscatter data for intelligent control of AUV science instrument carried out at MBARI. The objective of the study was to investigate techniques which use machine intelligence to enable event-response scenarios in the ocean sciences. The Gulper (Figure 1) is a water-sample acquisition system for an AUV. It is a syringe-like instrument which has a 2 liter bottle actuated by an electromagnetic trigger which can obtain the sample in about 2 seconds. The Gulper is mounted in the mid-section of an MBARI Bluefin-21 AUV with the intent to have it triggered by control software onboard the AUV. Our approach is to take historical data from a Hobilabs HS-2 HydroScat, recorded on earlier missions, and apply machine learning algorithms to learn classifications of the data. Scientists correlate the source data with the clustering results to indicate which clusters are of interest. The learned clusters are then encapsulated in a simple classifier onboard the AUV. During a mission, the AUV equipped with the classifier can process data from its HS-2 HydroScat sensor, normalizing the data appropriately before feeding it to the classifier algorithm to be categorized. If the data falls into one of the clusters previously marked by scientists as being of interest, the AUV controller linked to the online classifier algorithm will trigger the sample-acquisition by the Gulper.

The Gulper device was originally designed to be triggered by an *ad hoc* sequence initiated by the AUV control software. Our objective in the near-term is to enable the triggering based on contextual environmental sensing. Therefore, our pilot study makes use of clustering techniques, from Artificial Intelligence, to detect a Nepheloid layer with suspended materials using optical backscattering techniques using the HydroScat HS-2 instrument mounted on the AUV. The techniques in our study are completely generic and could be applied equally to include other measurable phenomena such as absorption and nutrient levels.

The specific machine learning techniques we have used in

this study include two unsupervised clustering algorithms, intended to identify the important distinguishing characteristics of bodies of points within a data sample. The use of two different approaches allows us to exploit the strengths of each of the approaches to provide a robust identification of interesting phenomena in the water-column. Specifically, we use a self-organising map (SOM) to construct an accurate classification of HS-2 Hydrosat data and an Inductive Monitoring System (IMS) clusterer which is effective at identifying data that lies outside the learned range of original training data. Clustering techniques have been applied to hydrosat data in the past [6], [14], but our work represents the first attempt to use clustered data to support *in situ* control of a scientific instrument in the ocean sciences.

This paper is organized as follows:

- section II motivates the use of AUVs for water sampling,
- section III describes the clustering techniques we are using
- section IV discuss the choices needed to be made during the learning phase for extraction of interesting feature sets
- section V shows the results from our classifiers for AUV transects
- section VI concludes and discusses future prospects

II. MOTIVATION

An environmental feature of great interest encountered very frequently in AUV surveys of the Monterey Bay region has been the mobilization and transport of shelf sediments as intermediate *nepheloid layers* (INLs), layers of turbid water containing constituents of shelf sediments. These layers are of scientific interest because their constituents and transports can have highly significant impacts on the coastal marine ecosystem. One of the important constituents is iron, which is enriched in shelf sediments. Iron acts as a limiting nutrient for the phytoplankton whose productivity fuels the oceanic food web [9]. It has also been shown that iron and other trace metals can strongly influence the toxicity of some species of phytoplankton [13], thus transport of these metals with INLs can influence the occurrence of harmful algal blooms. Another important constituent is a biological phenomenon — cysts of phytoplankton species. These cysts are resting stages that form when the plankton are under physiological stress, and they endure in shelf sediments for extended periods. Many harmful algal species form cysts. Transport of cysts into a growth-favorable environment in the shallow water column can initiate phytoplankton blooms. A third important constituent is carbon contained in “marine snow”, decaying organic matter from the pelagic ecosystem. Transport of this carbon off the continental shelf is an important part of the global ocean carbon budget. This budget ultimately determines the exchange of carbon between the ocean and atmosphere and hence has significant implications for understanding uptake of anthropogenic CO₂ by the ocean. INLs are not only scientifically important, but also methodologically tractable for applying machine learning to advance ocean science. INLs are clearly distinguished by their combination of high optical backscatter and low

chlorophyll fluorescence. This distinction is used in this study targeting INLs.

III. AUTOMATED CLASSIFICATION

In the Machine Learning community the idea of automating the partitioning of data, including signals, into classes is well-established [8]. Machine learning techniques used for constructing classifiers include a wide range of techniques appropriate in different contexts. In all approaches, the machine learning algorithm is presented with a large set of training data from which to learn. In this brief discussion we will concentrate on machine learning of classifiers, which are functions that can be used to classify unseen data into the categories learned from the training data set. An important distinction is between supervised and unsupervised techniques. In supervised techniques, the classifier is built by training it on labeled data sets, where the intended classification for each datum is provided as input. It is expected that these labels are generally accurate and they define the set of classes that the classifier is expected to learn. Following training the classifier is required to distinguish previously unseen positive and negative examples of the classes. Examples of supervised techniques include artificial neural networks and decision tree learners such as C5.0 [15]. In unsupervised learning the data is not labeled. The user need not have any a priori expectations of the structure of the classifications. The classifier partitions the data according to similarities and differences in the data itself. Subsequent classifications are then made by matching an unseen pattern to the class that best characterizes it. One of the benefits of unsupervised classification is that it does not require any prior knowledge about the structure of the data. The results of unsupervised classification have to be evaluated by use of cross-validation to an independent source of judgment. Unsupervised techniques include clustering, which build classifiers from patterns in the data itself. There are several approaches to clustering, including hierarchical approaches [10], which iteratively refine large clusters into smaller ones, and partitioning approaches, such as k-means [12]) which attempt to learn a single flat set of clusters.

All clusterers require that the data provided to them be presented as *feature vectors*: tuples of data values that characterize a single input datum point. The choice of features is critical in determining the effectiveness of the clustering process. If the selected features are not causally relevant to the association between the phenomena being observed and the classification of those phenomena, they are unlikely to disambiguate data sets. Much larger data sets will then be required to allow the clusterer to recognize the irrelevance of the features. On the other hand, if the feature set lacks causally relevant data then the clusterer will learn poorer quality classifications with poor discrimination and accuracy.

Although clustering techniques are considered unsupervised, different forms of clustering algorithms require more or less advice from the user about the structure of the expected classification. The most important piece of advice is the

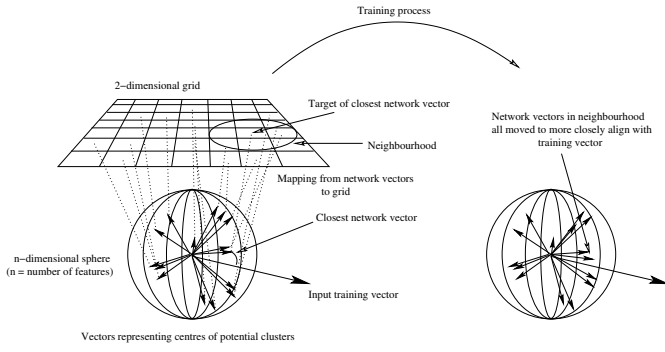


Fig. 2. The training process in a self-organizing map. At each iteration, the next input vector is presented to the network, the closest network vector identified and then it and its neighbors in the grid are moved closer to the input vector.

number of clusters. Determining this number is sometimes straightforward, being an obvious characteristic of the original data source, but often it is very difficult to provide and, where clusterers require it their application can involve careful analysis by hand of their performance with different settings. Clusterers that need to know the number of classes include k-means clusterers [4], [12], while quality threshold clustering [5] needs a different piece of advice, namely the minimum degree of similarity that will lead to data items being included in the same class. In this work we have used two clustering strategies: Self-Organizing Maps (SOM) or Kohonen Networks [11], and the method used in the Inductive Monitoring System (IMS) [7].

1) *Self-Organizing Maps*: The first classification approach that we used is closely related to the Kohonen Network classification technique [11], but is a variant of our own design [3]. This is a clustering approach that requires no advice on the number or size of clusters. It takes high-dimensional data vectors and projects them onto a lower-dimensional space.

In our application we used a 2-dimensional space, because a grid defines a simple neighborhood structure that can be efficiently updated. In our approach the grid is initialized with unit vectors obtained by normalizing randomly generated vectors of the same dimensionality as the input vectors. The projection of a given input vector is then performed by taking the dot product of the input vector with each of the unit vectors and then perturbing the closest unit vector to increase its similarity to the input vector (see Figure 2). The input vector is then discarded. The complete collection of input vectors is exposed to the network in this way until the network is not evolving anymore, and the resulting network forms the basis of the classification *landscape*. An outline of the algorithm is shown in pseudocode in Figure 3.

The numbers of training vectors that then associate most closely with each network vector can then be determined, leading to a 3-dimensional landscape over the 2-dimensional grid in which the tallest peaks represent the most common data patterns. Plateau and very small peaks often represent noise and which we remove by filtering the landscape to identify the

```

SOM(dataSet)
1  RANDOMLY INITIALIZE( $n \times n$  array of unit vectors,  $N$ )
2  while  $N$  not stable
3  do REDUCE NEIGHBOURHOOD SIZE
4    for each vector in dataSet
5    do  $best \leftarrow \text{CLOSESTVECTOR}(\text{vector}, N)$ 
6      for each  $v$  in  $\text{NBD}(best) \cup \{best\}$ 
7      do  $\text{REALIGNTOWARDS}(v, \text{vector})$ 
8  for each vector in dataSet
9  do  $best \leftarrow \text{CLOSESTVECTOR}(\text{vector}, N)$ 
10  INCREMENTCOUNT( $best$ )
11  for each vector in  $N$ 
12  do if  $\text{COUNTFOR}(\text{vector}) > \text{MAXCOUNT}(\text{NBD}(\text{vector}))$ 
13    then  $\text{ADDTO}(\text{clusterSet}, \text{vector})$ 
14  return clusterSet

```

Fig. 3. The SOM clustering algorithm

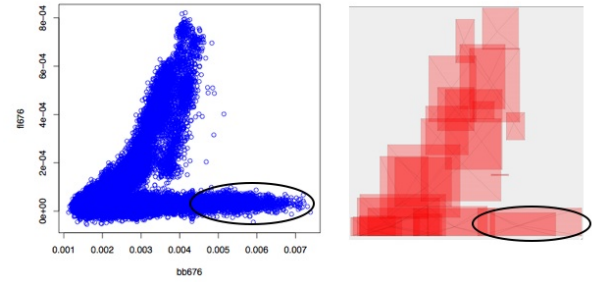


Fig. 4. An example of IMS clustering. On the left is input data of a nepheloid layer; on the right is the representation of the associated clusters found by IMS. The circular marked area corresponds to the feature of interest.

most information-rich clusters. This is achieved using a hill-climbing algorithm that traverses the landscape and merges small peaks with their tallest neighboring peaks. Although this loses some discriminatory power, it effectively restricts the set of clusters to those most likely to be of interest in a given application. While this hill-climbing step is not typical of SOM classifiers, our evaluation of the technique has demonstrated that the algorithm performs extremely well across a wide range of different types of data.

A cluster set can be evaluated by seeing how well it distinguishes patterns in a dataset with known characteristics that was not used in the classification phase. Having learned a cluster set it can be used to classify previously unseen data patterns, including as an online process. As soon as sufficient readings have been collected for a treated vector to be constructed, the new vector can be normalized against the training set and then compared against the learned cluster set to determine its characteristics. This is the approach we have taken in the experiment described in this paper.

2) *IMS clustering*: The IMS system has been used for trending analysis for anomaly detection for NASA's Space Shuttle Main Engine [7]. The key idea is to automatically extract a model from nominal data sets and then use this model

to monitor the system evolution for detecting divergence with the learned model.

To do so IMS uses a clustering algorithm to group sets of consistent parameter values found in training data. The inputs of the program are the learning data set represented as vectors and a cluster growth parameter ϵ . A small ϵ will produce a large number of clusters tightly connected to the data; a larger value for ϵ will produce fewer clusters but with a less accurate model of data. The choice of the ϵ provides a way to balance between accuracy (limiting the false positive output rate) of the model and knowledge size to be able to make the classification in real-time.

Using these parameters the algorithm will build a knowledge base containing clusters of related value ranges. Therefore entries of High and Low values in a cluster can be seen as defining the minimum bounding hyper-cube for the data captured by this cluster. An example of such data vectors is shown in Figure 5. An example of data clusters from a nepheloid layer is as shown in Figure 4.

bb 470	bb 676	chlor. fl.	Temp. (°C)
5.07×10^{-3}	3.65×10^{-3}	2.17×10^{-3}	12.1

	bb 470	bb 676	chlor. fl.	Temp.
High	6.02×10^{-3}	3.8×10^{-3}	1.02×10^{-3}	12.3
Low	4.07×10^{-3}	2.96×10^{-3}	2.2×10^{-3}	11.9

Fig. 5. Example of a sample vector (top) and a sample IMS cluster which includes the above vector in 4-dimensional space

The IMS algorithm is described in Figure 6. The training phase is starting with an empty set of cluster. We first normalize all the vectors in training data such that all attribute values are in the interval $[0, 1)$. For each data we try to find the “closest” cluster. This distance can be measured using various metrics (e.g euclidean or scalar). If no cluster exists or if the distance to the vector is greater than ϵ , a new cluster is created, else the input vector is binned into the closest cluster. We then merge redundant clusters which are found by growing the largest clusters by ϵ . This is done to avoid an explosion in the number of clusters which are identifying similar features.

The output of this algorithm is a set of hyper-cubes (or clusters) with their union containing all the vectors of the training data. We call this union the **envelope** (illustrated by the red area in Figure 4) of the model and its accuracy depends on the value of ϵ .

IV. EXPERIMENTAL SETUP

MBARI has a very large collection of HS-2 HydroScat data collected as part of routine AUV time-series missions in the Monterey Bay. HS-2 data consists of measurements of backscattering at two wavelengths and chlorophyll fluorescence at a single wavelength, together with depth. In our data, backscattering is measured at 470nm and 676nm, and fluorescence at 676nm (a chlorophyll wavelength). *High backscatter readings, especially combined with low fluorescence*, suggest

```

IMS(dataSet, epsilon)
1  clusterSet  $\leftarrow \emptyset$ 
2  NORMALIZE(dataSet)
3  for each vector in dataSet
4  do best  $\leftarrow$  CLOSESTCLUSTER(vector, clusterSet)
5     if best =  $\emptyset$  or DISTANCE(best, vector) > epsilon
6     then new  $\leftarrow$  CREATECLUSTER(vector)
7         ADD(new to clusterSet)
8     else GROW(best using vector)
9         MERGECLUSTERS(clusterSet)
10 return clusterSet

```

Fig. 6. The IMS clustering algorithm

a density of non-fluorescing material, such as occurring during an upwelling event.

We used data collected on separate AUV missions performed in 2003 and 2004; our clustering techniques applied to HS-2 data alone. We then compared this with the layer structure of the region of the bay in which the missions took place.

Clustering algorithms perform differently when exposed to a range of features. For example, the SOM algorithm works well when the dimensionality of input vectors is high while IMS might be impacted by high dimension data in tuning the ϵ threshold value. So for instance if the data is sparse, the ϵ threshold has to be bigger but, if the data is too dense a large value for ϵ will result in overly large clusters resulting in loss of disambiguation [1]. For this reason, our two clustering algorithms have two different feature sets reflecting the importance of dimensionality.

a) *SOM algorithm:* In the data we analyzed, in addition to the three basic HS-2 readings (bb470, bb676 and *fl* the chlorophyll fluorescence), the depth at which the readings were taken is recorded. Raw depth is not a good value to use as a feature since using it can lead to clusters that are sensitized to the depths at which training data is supplied. The consequence is that subsequent use of the learned clusters will be dominated by new depth values. We therefore did not use the depth values as features directly, but instead we moderated the other features by using the depth information to calculate extra features of possible relevance in the data set. In this study we used vectors of eight features:

$$\begin{aligned}
&< bb470, bb676, fl, bb470 \times d, bb676 \times d, fl \times d, \\
&bb470 + bb676, bb470/bb676 >
\end{aligned}$$

where d is depth.

In deriving the 4th, 5th and 6th features we assumed that the amount of backscatter and fluorescence would, in the absence of significant phenomena, be inversely proportional to depth, so that multiplying each by depth supports identification of interesting patterns in the data. In deriving the 7th and 8th features we made the further assumption that the amounts of backscatter at the two wavelengths tend to vary together, so

we added the last two features to offer an opportunity for clustering to exploit the possibility.

b) *IMS algorithm*: For this algorithm we have chosen to take a simpler input vector consisting of the following features:

$$\langle bb470, bb676, fl, temp \rangle$$

where *temp* is the temperature of the water. This feature set enabled scientists to directly interpret the clustering results with the raw data for purposes of validation while retaining its simplicity.

Since we were looking for features with high backscattering and low chlorophyll fluorescence, the *bbs* and *fl* provided the appropriate features correlating the data with the clustering results. We added temperature since nepheloid layers appear to be correlated to it, in addition to allowing us to differentiate results from the two algorithms.

V. EXPERIMENTAL RESULTS

We trained our system using the SOM and IMS algorithms on data from AUV runs in 2003 and 2004 with approximately 20,000 HS-2 readings. Our objective in this study was to determine if such clustering techniques could identify a specific feature of interest, namely high backscatter and low chlorophyll fluorescence.

In the case of SOM, this learning set was used to train a 30×30 network generating 14 clusters of which we identified one to be directly correlated with the feature of interest. Experimenting with differing ϵ parameters for the IMS algorithm using the same data sets, produced different sets of clusters. We retained two of them :

- $\epsilon = 0.08$ produced a set of 35 clusters giving a very precise envelope of the learning data.
- $\epsilon = 0.12$ extracted only 8 clusters with a looser envelope.

During the course of December 2006 through February 2007, these learned clusters (one for SOM and two for IMS using the above ϵ values) were installed on MBARI's AUV for four runs in the Monterey Bay as well as simulations on our desktops using data from 2004. Classification of data onboard the AUV was done at a frequency of 4 Hz. No Gulper was attached since in these runs we wanted to demonstrate the applicability of such classification techniques to sampling. The results of the classification for these learned clusters were stored in a log file which were then analyzed on shore. Figure 7 refers to one such run for SOM with the path of the AUV superimposed on a sampled region of the Bay. The yo-yo pattern marked in black indicates the detection of the desired feature of interest with high backscatter and low fluorescence and directly correlates to one of the learned clusters in SOM. No such clearly disambiguated data was available for IMS. Results showed us that IMS is less efficient in extracting interesting features from the samples especially when these features are close to each other.

The IMS clusters on the other hand performed well in spotting outliers which is consistent with its design philosophy. This was particularly evident in a subsequent run, when the AUV temporarily hit the sea-floor bottom. This event

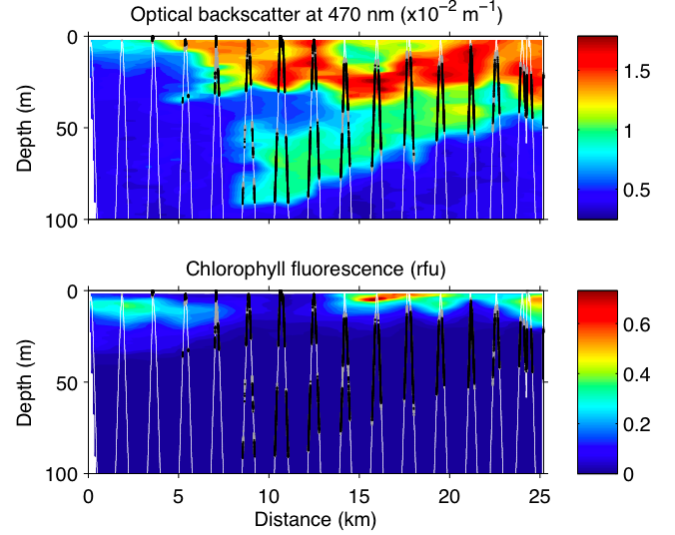


Fig. 7. Results of classification of the feature of interest, high backscatter and low chlorophyll uorescence, overlaid on gridded sections of the original oceanographic data. The yo-yo pattern shows the sampling locations; black points indicate feature detection with high confidence; grey points indicate feature detection with low confidence (along feature boundaries).

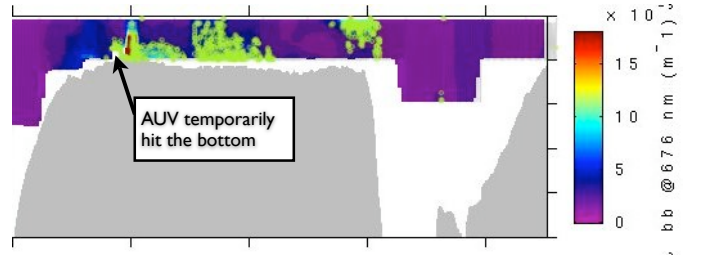


Fig. 8. Model outliers detected by IMS as shown in the light green areas

perturbed the sensor inputs. Figure 8 shows clearly where the IMS-based classifier identified data as outside the learned model. This run in addition to another where the HS-2 was inoperable, has led us to believe that IMS based techniques can be exploited in real-time to filter such outliers. Our next steps therefore will be to consider a hybrid version of these two algorithms using SOM for classification with IMS providing the relevant filtering for incoming data.

VI. CONCLUSIONS AND FUTURE WORK

Our preliminary results demonstrate that clustering can identify structure in the data set generated by HS-2 readings. Our initial experiments also show that the clusters, generated entirely unsupervised, offer a good correlation with the structure of relevant phenomena in the data set. SOM has been shown to be better in extracting the features from the data while IMS was able to easily point to data divergence. Such techniques could easily be generalized to other sensor data.

There are a number of other research areas that will benefit from enhancing the decision making capabilities of the AUV. This will apply not only to acquisition of water samples when and where events are detected, but also to modification of

AUV survey behavior. One example is the spread of larvae of invasive species in plumes from land-sea exchange. The ability to recognize a plume, remain within it, and sample throughout it as it is transported by coastal circulation will greatly advance understanding of the dispersion of these species. Another example relates to the extreme patchiness of plankton in the coastal ocean. Physical-biological interactions often create dense aggregations of plankton, and it is in these biological hot-spots that much of the activity of ocean life occurs. The most important patches to sample may be a small portion of the volume that the AUV surveys, thus rapid recognition and sampling of these features will greatly augment our understanding of ocean life. These examples represent a small slice of the spectrum of challenging research areas that will benefit from machine learning advancements applied to AUVs in coastal oceanography.

What such applications need is the precise contextual environment from which these samples are obtained. Existing approaches to obtaining such relevant samples have relied on a priori knowledge of the characteristics of the water column to ensure that a sampling device is triggered at the appropriate spatio-temporal scales. This has often resulted in the inability to view the entire life-cycle of the event or process being studied in the ocean or, which is worse, missing the event altogether due to the lack of instrument or platform presence, or in obtaining a sample that is not correlated with the intent of the scientific observation. In addition, where ship-based measurements are used to support such observation regimes, these face similar issues as well as steep costs. A capability is required for robotic devices to opportunistically trigger sampling devices (such as the Gulper) in the contextual environment in which scientific intent can be satisfied in addition to influencing the navigation and control of AUV's. Towards this end, our current research [16] is working to encapsulate deliberative techniques in Artificial Intelligence to enable real time decision making and adaptive sampling for ocean going platforms.

VII. ACKNOWLEDGMENTS

This work is partly funded by the David and Lucile Packard Foundation at MBARI. We would like to thank Hans Thomas and MBARI's Marine Operations AUV group.

REFERENCES

- [1] P. Berkhin. Survey of clustering data mining techniques. Technical report, Accrue Software, San Jose, CA, 2002.
- [2] F.P. Chavez, J. Ryan, S.E. Lluch-Cota, and M. Niquen. From anchovies to sardines and back: multidecadal change in the pacific ocean. *Science*, 299(5604):217–221, 2003.
- [3] M. Fox, M. Ghallab, G. Infantes, and D. Long. Robot Introspection through Learned Hidden Markov Models. *Artificial Intelligence*, 170(2):59–113, 2006.
- [4] J. A. Hartigan and M. A. Wong. A k -means clustering algorithm. *Applied Statistics*, 28(1):100–108, 1979.
- [5] L.J. Heyer, S. Kruglyak, and S. Yooseph. Exploring expression data: Identification and analysis of coexpressed genes. *Genome Research*, 9:1106–1115, 1999.
- [6] E.N. Hildebrand. The effect of particle size distribution on spectral backscattering coefficient. Master's thesis, Dalhousie University, 1999.
- [7] D. L. Iverson. Inductive system health monitoring. In *Int. conf. on Artificial Intelligence*, Las Vegas, NV, USA, jun 2004.
- [8] A. K. Jain, M. N. Murty, and P. J. Flynn. Data clustering: a review. *ACM Computing Surveys*, 31(3):264–323, 1999.
- [9] K.S. Johnson, F.P. Chavez, and G.E. Friederich. Continental-shelf sediment as a primary source of iron for coastal phytoplankton. *Nature*, 398(6729):697–700, 1999.
- [10] S.C. Johnson. Hierarchical clustering schemes. *Psychometrika*, 2:241–254, 1967.
- [11] T. Kohonen. *Self-Organisation and Associative Memory*. Springer Verlag, 1984.
- [12] J.B. MacQueen. Some methods for classification and analysis of multivariate observations. In *Proceedings of 5th Berkeley Symposium on Mathematical Statistics and Probability*, volume 1, pages 281–297, 1967.
- [13] M.T. Maldonado, M.P. Hughes, E.L. Rue, and M.L. Wells. The effect of fe and cu on growth and domoic acid production by pseudo-nitzschia multiseries and pseudo-nitzschia australis. *Limnol Oceanography*, 47(2):515–526, 2002.
- [14] M.A. Moline, R. Arnone, T. Bergmann, M.J. Oliver S. Glenn, C. Orrico, O. Schofield, and S. Tozzi. Variability in spectral backscatter estimated from satellites and its relation to in situ measurements in optically complex coastal waters. *Int. J. of Remote Sensing*, 25:1465–1468, 2004.
- [15] J.R. Quinlan. *C4.4: programs for machine learning*. Morgan Kaufmann, 1994.
- [16] K. Rajan, C. McGann, F. Py, and H. Thomas. Robust mission planning using deliberative autonomy for autonomous underwater vehicles. In *Intl. conf. on Robotics and Automation workshop on Robotics in challenging and hazardous environments*, 2007.
- [17] J.P. Ryan, H.M. Dierssen, R.M. Kudela, C.A. Scholin, K.S. Johnson, J.M. Sullivan, A.M. Fischer, E.V. Rienecker, P.R. McEnaney, and F.P. Chavez. Coastal ocean physics and red tides, an example from monterey bay, california. *Oceanography*, 18:246–255, 2005.