

Automated Learning and Discovery: State-Of-The-Art and Research Topics in a Rapidly Growing Field

Sebastian Thrun, Christos Faloutsos, Tom Mitchell, Larry Wasserman

Center for Automated Learning and Discovery
Carnegie Mellon University
Pittsburgh, PA 15213

September 1998

Abstract

This report summarizes the CONALD meeting, which took place June 11-13, 1998, at Carnegie Mellon University. CONALD brought together an interdisciplinary group of scientists, concerned with decision making based on data. This report is organized in two parts. The first part (pages 3-8) summarizes the CONALD meeting and highlights its main outcomes, beyond the individual workshop level. The second part (pages 9-30) summarize the results obtained in the individual workshops, discussing in depth promising research topics. This report is available through the Web at <http://www.cs.cmu.edu/~conald>.

CONALD was supported financially by the National Science Foundation under grant number IIS-9813354, which is gratefully acknowledged.

1 Introduction

The field of automated learning and discovery—often called data mining, machine learning, or advanced data analysis—is currently undergoing a revolution. The progressing computerization of professional and private life, paired with a sharp increase in memory, processing and networking capabilities of today’s computers, make it now more than ever possible to gather and analyze vast amounts of data. For the first time ever, the people all around the world are connected to each other electronically through the Internet, making available huge amounts of online data at a astonishingly increasing pace.

Sparked by these innovations, we are currently witnessing a rapid growth of a new industry, called the data mining industry. Companies and governments have begun to realize the power of computer-automated tools for systematically gathering and analyzing data. For example, medical institutions have begun to utilize data-driven decision tools for diagnostic and prognostic purposes; various financial companies have begun to analyze their customers’ behavior in order to maximize the effectiveness of marketing efforts; the Government now routinely applies data mining techniques to discover national threats and patterns of illegal activities in intelligence databases; and an increasing number of factories apply automatic learning methods to optimize process control. These examples illustrate the immense societal importance of the entire field.

At the same time, we are witnessing a healthy increase in research activities on issues related to automated learning and discovery. Recent research has led to revolutionary progress, both in the type methods that are available, and in the understanding of their characteristics. While the broad topic of automated learning and discovery is inherently cross-disciplinary in nature—it falls right into the intersection of disciplines like statistics, computer science, cognitive psychology, robotics, social sciences, and public policy—these fields have mostly studied this topic in isolation. So where is the field, and where is it going? What are the most promising research directions? What are opportunities of cross-cutting research, and what is worth pursuing?

2 The CONALD Meeting

To brainstorm about these and similar questions, Carnegie Mellon University (CMU) recently hosted the CONALD meeting. The goal of the meeting was to bring together leading scientists from the various disciplines involved, to brainstorm about the following central questions:

1. **State of the art.** What is the state of the art? What are examples of successful systems?
2. **Goals and impact.** What are the long-term goals of the field? What will be the most likely future impact of the area?
3. **Promising research topics.** What are examples of the most promising research topics that should be pursued in the next three to five years?
4. **Opportunities for cross-disciplinary research.** Which are the most significant opportunities for cross-cutting research?

The meeting, which took place in June 1998, drew approximately 250 participants. It featured seven plenary talks, given by Tom Dietterich (Oregon State University), Stuart Geman (Brown University), David Heckerman (Microsoft Research), Michael Jordan (MIT, now UC Berkeley), Daryl Pregibon (AT&T Research), Herb Simon (CMU), and Robert Tibshirani (Univ. of Toronto, now Stanford Univ.). Apart from the plenary talks, which were aimed at familiarizing researchers from different scientific community with each other’s research, the meeting was basically a collection of seven workshops

1. **Learning Causal Bayesian Networks** (see page 9)
organized by Richard Scheines and Larry Wasserman.
2. **Mixed-Media Databases** (see page 12)
organized by Shumeet Baluja, Christos Faloutsos, Alex Hauptmann and Michael Witbrock.
3. **Machine Learning and Reinforcement Learning for Manufacturing** (see page 14)
organized by Sridhar Mahadevan and Andrew Moore.
4. **Large-Scale Consumer Databases** (see page 18)
organized by Mike Meyer, Teddy Seidenfeld, and Kannan Srinivasan.
5. **Visual Methods for the Study of Massive Data Sets** (see page 19)
organized by Bill Eddy and Steve Eick.
6. **Learning from Text and the Web** (see page 22)
organized by Yiming Yang, Jaime Carbonell, Steve Fienberg, and Tom Mitchell.
7. **Robot Exploration and Learning** (see page 26)
organized by Howie Choset, Maja Matarić, and Sebastian Thrun

Table 1: CONALD workshop topics. A detailed report of each workshop can be found in Part 2 of this report.

(see Table 1), where workshop participants discussed a specific topic in depth. Each workshop was organized by an inter-disciplinary team of researchers, and the topic of the workshops infringed on research done in several areas. At the last day, all workshop participants met in a single room for two sessions called “thesis topics,” where workshop chairs summarized the results of their workshop and laid out concrete, promising research topics, as examples of feasible and promising topics for future research.

CONALD was attended by approximately 240 researchers, the majority of whom were computer scientists or statisticians. Apart from the invited plenary speakers, each workshop invited up to two leading scientists, using funds provided by the National Science Foundation. The workshop served two primary purposes. Apart from the agenda to brainstorm and identify promising research directions, it also brought together leading scientists from areas such as artificial intelligence, databases, and statistics, and initiated a dialogue which we hope will carry on in the future.

3 The Need For Cross-Disciplinary Research

A key objective of the CONALD meeting was to investigate into the role of a cross-disciplinary approach. Historically, issues of automated learning and discovery have been studied by various scientific disciplines, such as statistics, computer science, cognitive psychology, social sciences, and public policy. In many cases, each discipline pursued its research in isolation, studying specific facets of the general problem, and developing a unique set of methods, theory, and terminology. To illustrate this point, Table 2 shows a modified version of a “statistics-to-AI dictionary,” shown by Rob Tibshirani in his plenary talk (and later augmented by Andrew W. Moore) to illustrate difference in

Statistics Term	AI Term
statistics	data learning
regression	progression, straight neural network
discrimination	pattern recognition
prediction sum squares	generalization ability
fitting	learning
empirical error	training set error
sample	training set
experimental design	active learning

Table 2: An modified version of Tibshirani’s Statistics-to-Artificial Intelligence dictionary illustrates the differences in terms used in two scientific field concerned with about the same questions.

terminology.

There was a broad consensus that the issues at stake are highly interdisciplinary. Workshop participants and organizers alike expressed that each discipline has studied unique aspects of the problem and therefore can contribute a unique collection of approaches. Statistics—undoubtedly the field with the longest-reaching history—has developed powerful methods for gathering, learning from and reasoning with data, often studied in highly restrictive settings. Researchers in AI have explored learning from huge datasets with high-dimensional feature spaces (e.g., learning from text). Database researchers have devised efficient method for storing and processing huge datasets, and they have devised highly efficient methods for answering certain types of questions (such as membership queries). Various applied disciplines have contributed specific problem settings and datasets of societal importance. Many participants expressed that by bringing together these various disciplines, there is an opportunity to integrate each other’s insides and methodologies, to gain the best of all worlds. In addition, an interdisciplinary discourse is likely to reduce the danger of wasting resources by re-discovering each other’s results.

4 State Of The Art, Promising Research Directions

Characterizing the state-of-the-art is not an easy endeavor, as the space of commercially available approaches is large, and research prototypes exist for virtually any problem in the area of automated learning and discovery. Thus, we will attempt to give some broad characterizations that, in our opinion, most people agreed to.

There was an agreement that function fitting (which often goes by the name of supervised learning, pattern recognition, regression, approximation, interpolation—not all of which mean the exact same thing) appears now to be a well-understood problem. This is specifically the case when feature spaces are low-dimensional and sufficient data are available. There exists now a large collection of popular and well-understood function fitting algorithms, such as splines, logistic regression, Back-Propagation, decision trees, and so on. Many of these tools form the backbone of commercial data-mining tools, where they analyze data and predict future trends.

Clustering of data in low-dimensional spaces has also been studied extensively, and today we possess a large collection of methods for clustering data, which are specifically applicable if feature spaces are low-dimensional and plenty data is available.

Of course, data mining is more than just applying a learning algorithm to a set of data. Existing tools provide powerful mechanisms for data preparation, visualization, and interpretation of the results. Many of the tools work well in low-dimensional, numerical feature spaces—yet they cease to work if the data is high-dimensional and non-numerical (such as text). There was a reasonable consensus that such work is important, and better methodologies are needed to data preparation, processing, and visualization.

The workshop sessions generated a large number of research topics. Despite the fact that the workshops were organized around different problem/application domains, many of these topics co-occurred in multiple workshops. Among the most notable were the following:

- **Active learning/experimental design.** Active learning (AI jargon) and experimental design (statistics jargon) addresses the problem of choosing which experiment to run during learning. It assumes that during learning, there is an opportunity to influence the data collection. For example, a financial institution worried about customer retention might be interested *why* customers discontinue their business with the institution, so that potential candidates can be identified and the appropriate actions can be taken. It is impossible, however, to interview millions of customers, and in particular those who changed to another provider are difficult to ask—so whom should one call to learn the most useful model? In robot learning, to name a second example, the problem of “exploration” is a major problem. Robot hardware is slow; yet, most learning methods depend crucially on a wise choice of learning data. Active learning addresses the question of how to explore.
- **Cumulative learning.** Many practical learning problems are characterized by a continual feed of data. For example, databases of customer transaction or medical records grow incrementally. Often, the sheer complexity of the data and/or the statistical algorithm used for their analysis prohibits evaluating the data from scratch every day. Instead, data has to be analyzed cumulatively, as they arrive. This problem is specifically difficult if the laws underlying the data generation may change in non-obvious ways: For example, customers’ behavior can be influenced by a new regulation, a fashion, a weather pattern, a product launched by a competitor, a recession, or a scientific discovery. Can we devise cumulative learning algorithms that can incrementally incorporate new data, and that can adapt to changes in the process that generated the data?
- **Multitask learning.** Many domains are characterized by families of highly related (though not identical) learning problems. Medical domains are of this type. While each disease poses an individual learning task for which dedicated databases exist, many diseases share similar physical causes and symptoms, making it promising to transfer knowledge across multiple learning tasks. Similar issues arise in user modeling, where knowledge may be transferred across individual users, and in financial domains, where knowledge of one stock might help predicting the future value of another. Can we devise effective multi-task learning algorithms, which generalize more accurately through transferring knowledge across learning tasks?
- **Learning from labeled and unlabeled data.** In many application domains, it is not the data that is expensive; instead, obtaining labels for the data is a difficult and expensive process. For example, software agents that adaptively filter on-line news articles can easily access vast amounts of data almost for free; however, having a user label excessive amounts of data (e.g., expressing his level of interest) is usually prohibitive. Can we devise learning algorithms that exploit the unlabeled data when learning a new concept? If so, what is the relative value of labeled data compared to unlabeled data?

- **Relational learning.** In many learning problems, instances are not described by a static set of features. For example, when finding patterns in intelligence databases, the relation between entities (companies, people) is of crucial importance. Entities in intelligence databases are people, organizations, companies and countries; and the relation of them is of crucial importance when finding patterns of criminal activities such as money laundering. Most of today's learning algorithms require fixed feature vectors for learning. Can we devise relational learning algorithm that consider the relation of multiple instances when making decisions?
- **Learning from extremely large datasets.** Many dataset are too large to be read by a computer more than a few times. For example, many grocery stores collect data of each transaction, often producing gigabytes of data every day. This makes it impossible to apply algorithms that require many passes through the data. Other databases, such as the Web, are too large and too dynamic to permit exhaustive access. Many of the existing algorithms exhibit poor scaling abilities when the dataset is huge. Can we devise learning algorithms that scale up to extremely large databases?
- **Learning from extremely small datasets.** At the other extreme, there are often databases that are too small for current learning methods. For example, in face recognition problems, there is often just a single image of a person available, making it difficult to identify this person automatically in other images. In robotics, the number of examples is often extremely limited; yet many of the popular learning algorithms (e.g., genetic programming, reinforcement learning) require huge amounts of data. Obviously, the key of learning from scarce data lies in prior knowledge (see below)? What other ways are there to learn from scarce datasets?
- **Learning with prior knowledge.** In many cases, substantial prior knowledge is available about the phenomenon that is being learned. For example, one might possess knowledge about political events, laws and regulations, political events, personal preferences, and economics essential to the prediction of exchange rates between currencies. How can we incorporate such knowledge into our statistical methods? Can we find flexible schemes that facilitate the insertion of diverse, abstract, or uncertain prior knowledge?
- **Learning from mixed media data.** Many data sets contain more than just a single type of data. For example, medical datasets often contain numerical data (e.g., test results), images (e.g., X-rays), nominal data (e.g., person smokes/does not smoke), acoustic data (e.g., the recording of a doctor's voice), and so on. Existing algorithms can usually only cope with a single type data. How can we design methods that can integrate data from multiple modalities? Is it better to apply separate learning algorithms to each data modality, and to integrate their results, or do we need algorithms that can handle multiple data modalities on a feature-level?
- **Learning casual relationships.** Most existing learning algorithms detect only correlations, but are unable to model causality and hence fail to predict the effect of external controls. For example, a statistical algorithm might detect a strong correlation between chances to develop lung cancer, and the observation that a person has yellow fingers. What is more difficult to detect is that both, lung cancer and yellow fingers are caused by a hidden effect (smoking), and hence providing alternative means to reduce the yellowness of the fingers (e.g., a better soap) is unlikely to change the odds of a person developing cancer—despite the correlation! Can we devise learning algorithms that discover causality? If so, what type assumptions have to be made, to extract causality from a purely observational database, and what implications do they have?

- **Visualization and interactive data mining.** In many applications, data mining is an interactive process, which involves automated data analysis and control decisions by an expert of the domain. For example, patterns in many large-scale consumer or medical databases are often discovered interactively, by a human expert looking at the data, rearranging it, and using computer tools to search for specific patterns. Data visualization is specifically difficult when data is high-dimensional, specifically when it involves non-numerical data such as text. How can we visualize large and high-dimensional data sets? How can we design interactive tools that best integrate the computational power of computers with the knowledge of the person operating it?

This list has been compiled based on the outcomes of the individual workshops, the invited talks, and various discussions that occurred in the context of CONALD. Virtually all of these topics are cross-disciplinary in nature. While this list is necessarily incomplete, it covers the most prominent research issues discussed at this meeting.

Enclosed to this report are the individual workshop summaries, written by the individual workshop organizers with the help of volunteer scribes.

Acknowledgements

CONALD was financially supported by the National Science Foundation (NSF), which is gratefully acknowledged. The meeting was also sponsored by CMU's Center for Automated Learning and Discovery, created by researchers in Computer Science and Statistics in 1997. Finally, we would like to thank all participants who contributed to a lively meeting, and who were of course essential.

Workshop 1: Learning Causal Bayesian Networks

Organizers: Richard Scheines and Larry Wasserman

Scribes: Stella Salvatierra and Kary Myers

Over the last decade, several research groups in computer science, philosophy, and statistics have made substantial progress on connecting causal hypotheses represented as directed acyclic graphs (DAGs) to a simple but broad class of statistical models called Bayesian networks, on specifying algorithms to learn about causal hypotheses from background knowledge and statistical data, and on characterizing the limits of what can and cannot be learned from statistical data under a variety of assumptions. In this Workshop, six groups presented work at the state of this art, and in each case identified several research topics that would constitute excellent Ph.D. topics. The papers presented are all available on the web at:

<http://www.andrew.cmu.edu/course/80-512/conaldpapers/papers.html>

Phil Dawid, a statistician from University College in London, began the workshop by examining the semantics of DAGs that are interpreted causally. Dawid showed how DAG 1: $X=AEY$ was semantically distinguishable from DAG 2: $Y=AEX$. Even though neither graph puts any constraint on the joint probability distribution over X, Y , the DAGs make very different claims about the result of interventions. Dawid made clear the extra structure we must introduce in order to distinguish between an Intervention DAG and a regular DAG, where the former is capable of expressing causal hypotheses but the latter is not. He then proceeded to distinguish between two alternative interpretations of the intervention DAG: the structural interpretation and the probabilistic interpretation. In the former, one assumes that each variable is a deterministic function of all of its causes, some of which are not measured. Making assumptions about the parameteric form of these functions allows us, in certain circumstances, to make inferences about the values of these unmeasured causes. Quite different, and usually weaker, inferences are supported by the probabilistic interpretation, and Dawid's prime concern was to invoke Occam and prevent the community from embracing unnecessary theoretical entities. A research topic that arose from this work is as follows.

Topic I: What calculations can be performed with the structural interpretation of Causal DAGs that cannot be performed with the probabilistic interpretation, and vice versa? Further, what sorts of empirical evidence, especially the kind typically collected by epidemiologists and social scientists, would support a structural interpretation?

The second paper represents collaboration by an epidemiologist (Jamie Robins from Harvard), a statistician (Larry Wasserman from CMU), and two philosophers (Peter Spirtes and Richard Scheines from CMU). This group explored the limits of what can be learned about causal hypotheses from finite samples. Most of the results in the field concern the asymptotic behavior of algorithms, that is, the behavior when the sample size approaches infinity. Robins, Scheines, Spirtes and Wasserman investigated the types of statistical convergence one can achieve in distinguishing causal hypotheses (as opposed to distinguishing probabilistic hypotheses) when unmeasured confounders cannot be ruled out a priori. Interestingly, even assuming faithfulness it turns out that pointwise but not uniform consistency is usually all that can be obtained. This means that on finite samples, we cannot calculate confidence intervals around specific causal parameters of the sort statisticians usually derive. This leads to:

Topic II: What sort of prior information would one need about the amount of unmeasured confounding in order to produce usefully narrow confidence intervals around causal effects? In particular, doing a sensitivity analysis of the behavior of algorithms on finite samples could serve to illuminate how serious a limitation pointwise vs. uniform consistency imposes in this setting.

The third paper came from a team of computer scientists at Berkeley: Nir Friedman, Kevin Murphy, and Stuart Russell, and involved the prospects for handling dynamic Bayesian networks, that is Bayesian networks that model time. Friedman, Murphy, and Russell presented an algorithm called Structural EM, which combines parameter estimation with model search, and they applied it to dynamic Bayesian networks and showed how to extend the technique to dynamic models with unmeasured variables. Their work leads to:

Topic III: When should we introduce unmeasured variables into dynamic Bayesian networks? How can we score dynamic Bayesian networks (with unmeasured variables) efficiently enough to search a large space of models in feasible time?

The fourth paper came from the Intelligent Systems Program and Center for Biomedical Informatics at the University of Pittsburgh, and involved automatically discretizing continuous variables. The technology available for handling Bayesian networks usually restricts one to either all continuous or all discrete variables, but most real data sets involve both. Stefano Monti and Greg Cooper described an algorithm to automatically discretize continuous variables under the assumption that the continuous measures were really discrete variables made continuous by noise. Since the main connection between causal structure and probability distributions is through conditional independence, the key question that arose out of this work is:

Topic IV: How can we discretize continuous variables and preserve conditional independence? Does Monti and Cooper's algorithm have this property?

The fifth paper was presented by Thomas Richardson, a statistician/philosopher at the University of Washington, Seattle. Richardson presented work motivated by the problems that arise because several causal models can be statistically indistinguishable. Richardson presented a class of graphical objects called Ancestral Graphs, which serve two purposes simultaneously. One, they represent an equivalence class of models that are indistinguishable by conditional independence relations, and include models with unmeasured variables (which makes the class infinite). Two, Ancestral Graphs capture features of the causal structure that are common to every member of an equivalence class, as well as features that are not. An important topic to arise from this research is:

Topic V: How can we search over equivalence classes instead of individual models? How can we parameterize Ancestral Graphs in such a way as to permit maximum likelihood estimation and statistical scoring, e.g., BIC, MML, etc?

The sixth paper was authored by Chris Meek and Dan Geiger, both of whom do research on Bayesian networks at Microsoft Research. Meek and Geiger are doing work crucial to score based model search, which involves specifying a class of models to be searched, e.g., DAGs with no latent variables, a score with which to evaluate each model, e.g., BIC, AIC, etc., and an algorithm for searching the space, e.g., greedy search. Almost all scores involve two terms, one that expresses the deviation between the data and the best prediction obtainable by the model, and another that penalizes model complexity. The problem is that for most latent variable models, the model's true complexity is

usually unknown. Geiger and Meek are working on the problem by reparameterizing causal models into "natural" parameters, and then investigating whether, for example, the natural parameter space is a linear exponential family, a curved exponential family, or a stratified exponential family. The answer can prove crucial in calculating the appropriate model complexity. Meek and Geiger also investigate a process called implicitization, which is available in computer algebra packages like Mathematica to solve for the complexity of the model.

Topic VI: Can implicitization or other techniques introduced by Geiger and Meek be automated and made sufficiently efficient to allow appropriate score based model search over discrete causal models that include latent variables?

Workshop 2: Mixed-Media Databases

Organizers: Shumeet Baluja, Christos Faloutsos, Alex Hauptmann and Michael Witbrock

Scribe: Sean Slattery

1 State of the Art

There is significant progress in the processing of a single medium (such as images, time sequences etc.). “Query by Example” seems to be the state of the art (‘find images with similar color distribution like a sunset image’). Successful multimedia processing systems include commercial systems (QBIC, Virage, Excalibur, musclefish) as well as university prototypes (e.g., Informedia, VideoQ, EigenFaces, Sphinx, Photobook).

2 Goals and Impact

It seems that no clear categorization or organization of the various research efforts concerning mixed-media databases exists. Moreover, currently no good evaluation tests exist to allow comparison of one system with another. One of the immediate goals of the field should be to map out what the various problems are (image and video segmentation, classification, feature extraction, information retrieval), how they relate to each other and how performance on each problem should be measured. The possibility of standard test corpora becoming available, similar to those used in the MUC and TREC evaluations, was discussed.

Additional long term goals considered were:

- What is the best way (from both a HCI and Mixed-Media viewpoint) to pose multimedia queries to a system? Can users use multimedia queries more effectively than existing text-only queries?
- How can a mixed-media database be processed to allow query by concept, rather than the query by “pixel” (ie., image statistics) that we currently use?
- Can we do Data Mining on Mixed Media databases? How can advances in Data Mining be used?
- How can we do information extraction from multimedia databases?
- How can we deal effectively with the fact that media objects may be “about” more than one thing? For example, a sunset photograph belongs to the class of “open-door scene”, as well as to the class of “romantic” scenes etc.

In terms of potential impact, the following fields could benefit from mixed-media databases:

- Medicine (mixed-media patient records, with images, electro-cardiograms, demographic data, and symptoms, each evolving over time)
- Entertainment (video processing, query by image or video content)
- Education (searching art collections, or reference material, by content)
- Industrial/Military Intelligence
- Finance (Forecasting; similarity detection; outlier detection)

3 Promising Research Topics

Broadly, the research topics proposed are as follows:

1. Scale up issues - How can we deal with more data? We clearly need to reduce the size. Given a data matrix (where, e.g., rows are objects and columns are features), we can reduce the size in two ways:
(a) Reducing the number of rows (e.g., through sampling, clustering etc.) or (b) Reducing the number of columns, with dimensionality reduction methods.
2. Cross-media Training/Learning: How can we obtain better performance on a task by leveraging information from another source? For example, we want to correlate speech segments to text; or speech segments to images, faces or specific persons. Another important area is to use motion.
3. Query by Concept: Currently, we can answer queries by color, shape, texture and motion. Thus, if the user gives a sample image of a soccer game (human-like blobs in a green background), current systems will find images with similar amounts of green etc. The goal is to also find images related to soccer, like the image of the captain of the winning team raising the world cup, even if the colors and shapes are completely different.
4. Summarization: Given a large collection of multimedia objects (say, video clips on the same topic), we would like to collapse them to a single video clip, eliminating duplicates or near duplicates. Also, given a time-stamped set of multimedia objects, we would like to detect and track topics.

4 Opportunities for Cross-Disciplinary Research

There are ample opportunities for cross-disciplinary work. Some of the potential fields that participate are the following, along with the tools they can offer:

- Digital Signal Processing: Absolutely necessary for the early processing of images, voice and video. Powerful techniques developed there include the Hidden Markov Models for voice processing, Wavelets and Fourier analysis for voice and images, Linear Predictive Coding for speech. The efforts on MPEG-4 seem very promising.
- Databases: A lot of methods, have been developed there, like R-trees and spatial access methods for handling n-d points and regions; fast clustering methods; Association Rules for data mining.
- AI and Machine Learning: Classification Trees, Boosting, Artificial Neural Networks, are some of the many methods that this discipline has to offer.
- Statistics: A large number of methods are valuable, like the Singular Value Decomposition (also known as Latent Semantic Indexing, Karhunen-Loeve transform, or Principal Component Analysis); techniques for regression, forecasting (ARIMA methodology); Probabilistic Image Models using self-similarities and fractals.
- Information Retrieval: Techniques like the Vector Space Model, Relevance Feedback, Clustering, are some of the time-tested methods that IR can offer for text processing.
- Parallelism/Hardware: Next-generation parallel hardware, like the Active/Intelligent Disks, seem very promising for processing large mixed-media databases.

Workshop 3: Reinforcement Learning and Machine Learning for Manufacturing

Organizers: Sridhar Mahadevan, Andrew Moore

Scribe: Joseph O’Sullivan

1 Introduction

In recent years, there has been a flurry of research on statistics and machine learning applied to decision making and control. Exciting progress has been made in many areas, including reinforcement learning, neural networks, and diagnostic Bayesian networks. Applications are emerging in the control of continuous processes, batch processes (e.g. wafer fabrication), probabilistic diagnosis, and industrial engineering tasks such as optimal control of transfer lines or production scheduling. The workshop drew together researchers from industrial labs and universities working on manufacturing problems and researchers in the fields of reinforcement learning, robotics, statistics and machine learning. The talks and discussions at the workshop covered a great deal of ground, and here we will restrict ourselves to the main themes.

2 State of the Art

One notable trend that emerged is the maturity of reinforcement learning in its applicability to actual industrial problems. A number of speakers stressed that reinforcement learning is close to real-world deployment. Schneider (CMU) [7] described a novel production scheduling system formulated as a Markov Decision Process (MDP) [5]. Mahadevan (MSU) showed that conventional control heuristics of multi-product transfer lines can be significantly outperformed by reinforcement learning, using a model-free average-reward reinforcement learning method for semi-Markov decision processes [3]. Tadepalli (Oregon State) [9] introduced another important manufacturing domain in which reinforcement learning appears likely to make a significant mark—the control of autonomous guided vehicles (AGVs) on factory floors. Riedmiller (Univ. of Karlsruhe) applied reinforcement learning to heating control of liquid tanks [6]. Finally, Schulte (CMU) innovatively used reinforcement learning to control lighting conditions in an intelligent workplace [8].

A warning note raised in the workshop is that machine learning studies (which are based on standard feature-vector datasets) are often idealized and far removed from realistic industrial practice. This issue was highlighted in a talk by Tjoelker (Boeing). In attempting to model causes of alarms there are several very different types of data available: Time series streams of sensor data (gigabytes, multiple frequencies, hundreds of sensors), occasional, rare, alarm events, process log data and human-defined measures of performance. Successful systems have to be able to deal with this variety of inputs and time scales. The importance of this concept was further reinforced by Hagan (Amsterdam)

The role of the learner extends down to the actual data gathering — what data should be gathered and what should be instrumented. Tjoelker mentioned the importance of this issue at Boeing, and Gür Ali described attempts at General Electric to mesh together the highly non-linear non-parametric worlds of Neural Networks with the formal, and highly useful Experiment Design literature from statistics. A discussion then ensued on when active gathering of data may be needed.

As described in a talk by Moore (CMU), some learning scenarios do not require us to try to learn

a complete model of the process dynamics. They simply need us to be expert about parts of the state space close to the optimal. The simplest case of this is the apparently humble task of parameter tweaking for noisy systems. Moore advocated the view that this form of active learning problem is of great importance whether the parameters being tweaked are for an algorithm, a real manufacturing process, a simulation, or a scientific experiment. Moore then introduced a new algorithm called Q2, designed to provide an approximation to "Autonomous Response Surface Methods" for experiment design [4].

Finally, two topics were discussed where the state of the art dramatically lags requirements: learning to control large scale distributed systems, and learning to control hierarchically organized systems.

The standard test-bed for large scale distributed systems appears to be packet-based telecommunication systems. Nowe (Free University) described this problem in detail, surveyed key earlier work applying reinforcement learning [1], and showed the importance of taking into account the actual costs of packets involved in learning. There was a spirited discussion of the observation that communications network routing is not an isolated case for reinforcement learning, and this, as yet, little-researched area could also be applicable to power-delivery networks, road traffic management, and ecologies.

Schneider, Mahadevan, Riedmiller each provided different perspectives on learning to control hierarchically organized systems. At the essence, with multiple systems (such as the various machines lying around a factory), it is essential that we avoid modelling the whole state-space monolithically, but that we should find a way to decompose the problem into a hierarchy of reinforcement learning systems. One promising approach, described in a keynote talk by Dietterich [2] is the MAX-Q architecture. Somewhat further in the future is the possibility of systems that learn their own decomposition.

3 Research Directions

A large number of detailed suggestions were proposed as promising topics for further investigation. Four broad avenues emerged, reflecting the key goals of the participants.

- to extend the range of domains to which we can apply current algorithms.
- to understand and interpret what these algorithms learn.
- to improve our usage of these algorithms.
- to continue inventing novel approaches or algorithm to tackle unaddressed needs.

How should we seek to extend the range of our current algorithms? First, by extending current discrete algorithms and convergence proofs to continuous domains. This direction implies developing better representations of continuous domains. Second, by extending current reinforcement learning theory in the discounted reward framework to the average reward framework. As discussed by Mahadevan and Tadepalli, average reward is a more natural metric for optimizing factories. Third, by extend current algorithms to work on optimization of time varying processes. Fourth, by applying a decomposition approach to optimizing large factory systems. Interesting methods, such as MAXQ [2], have indeed been proposed, but their scalability to real manufacturing domains is not clear. Such domains could be modeled as multi-agent reinforcement learning systems, or as a centralized optimizer, in either case with clear hierarchical elements.

Manufacturing practitioners would like the result of a learning system to be transparent models that can be easily visualized and understood. It is crucial that our learning algorithms allow this. Due

to the high cost of losing a shift, management is unwilling to accept black boxes that simply claim to improve processes. Instead, the goal should be a theoretical model of what factors influence overall performance, which can serve as a diagnostic aid to improve availability or robustness. Methods for cautious exploration need to be investigated, to ensure that controlled systems can provide guarantees of remaining in “safe” regions. We need to explore better methods for value function approximation. Current approximators are not robust or fast but most importantly, there is a mismatch between the theory, and between what is used in practice (almost every large scale study of reinforcement learning has used neural nets, despite theoretical counter-examples showing instability in weights).

We should seek to extend our methods to bootstrap or otherwise utilize the knowledge that practitioners have developed. In other words, develop methods that explicitly utilize what is known about the domain, e.g. physical laws in some domains, business rules in others. This can be extended in terms of more flexible ways of inserting this knowledge, coping when the constraints based on physics or from an expert are wrong, using this knowledge for diagnostics of the learned model, and generating better models. Alternatively, we should seek to incorporate a better understanding of the dynamics of the underlying system, instead of treating every optimization problem as a black box.

Most importantly, there is a pressing need for further innovation in algorithms, theory, and experiments. Further work is needed to create successful control systems which can relearn over a long lifetime, so that we don’t have to retrain from scratch each time a component of the model is altered. More sophisticated or explicit methods for dealing with uncertainty are also desirable. In particular, there is a desire for algorithms that will work well independent of stochastic inputs. Also, algorithms are needed which deal with probabilistic distributions on the outputs – the current fall-back of creating multiple state variables is not sufficient. We need to investigate speedup methods for accelerating convergence, including shaping, which biases the distribution of problems from relatively easy to more difficult. Also, constraints are needed that reduce sample complexity of reinforcement learning – optimal solutions can now be learned, but at prohibitive costs due to exploration.

Finally, some lines of research which have parallels in other fields should be explored — as examples, research by machine learners into active learning is paralleled by research from the statistical community into experimental design, and the needed research into utilizing domain knowledge is reflected in statistical bayesian frameworks.

4 Outlook

To summarize, manufacturing is an excellent domain for machine learning and reinforcement learning, since it is both in a relatively mature state allowing for state-of-the-art applications, but also presents some difficult real-world issues that require fundamental advances in the underlying framework and algorithms. Much of the state of the art presented in the workshop demonstrates maturity and is ripe for commercial development. While the lack of transparency hinders their effectiveness, these systems can be seen to outperform non-learning alternatives. We would like to thank the participants and to urge them to continue their current successes.

References

- [1] J. Boyan and M. Littman. Packet Routing in Dynamically Changing Networks: A Reinforcement Learning Approach. In *Neural Information Processing Systems (NIPS)*, 1993.

-
- [2] T. Dietterich. The MAXQ Method for Hierarchical Reinforcement Learning. In Jude Shavlik, editor, *International Conference on Machine Learning (to appear)*, 1998.
 - [3] S. Mahadevan, N. Marchalleck, T. Das, and A. Gosavi. Self-Improving Factory Simulation using Continuous-Time Average-Reward Reinforcement Learning. In *Proceedings of the 14th International Conference on Machine Learning (IMLC '97)*, Nashville , TN. Morgan Kaufmann, July 1997.
 - [4] A. W. Moore, J. Schneider, J. Boyan, and M. S. Lee. Q2: Memory-based active learning for optimizing noisy continuous functions. In Jude Shavlik, editor, *International Conference on Machine Learning (to appear)*, 1998.
 - [5] M. Puterman. *Markov Decision Processes: Discrete Dynamic Stochastic Programming*. John Wiley, 1994.
 - [6] M. Riedmiller. High Quality Thermostat Control by Reinforcement Learning – A Case Study. In *Conference on Automated Learning and Discovery*, 1998.
 - [7] J. G. Schneider, J. A. Boyan, and A. W. Moore. Value Function Based Production Scheduling. In Jude Shavlik, editor, *International Conference on Machine Learning (to appear)*, 1998.
 - [8] J. Schulte and S. Thrun. Reinforcement Learning for Intelligent Building Control. In *Conference on Automated Learning and Discovery*, 1998.
 - [9] P. Tadepalli. Average Reward Reinforcement Learning for Optimizing Factories. In *Conference on Automated Learning and Discovery*, 1998.

Workshop 4: Large-scale Consumer Databases

Organizers: Mike Meyer, Teddy Seidenfeld and Kannan Srinivasan

Scribe: Xu Fan

Large-scale consumer data bases provide the opportunity for novel approaches to marketing and customer services. For this workshop we invited submissions of work-in-progress from which the organizers selected three papers for advance distribution. These became the foci of the workshop's activities. In the order of their presentation, the following highlighted challenges of datamining consumer records based on cluster analysis:

1. "A formal statistical approach to collaborative filtering," Lyle H. Unger and Dean P. Foster, CIS Department, University of Pennsylvania
2. "Automated discovery of discriminant rules for a group of objects in databases," Tae-wan Ryu and Chris F. Eick, Department of Computer Science, University of Houston.
3. "Credit card attrition modeling," Peter Johnson and Jim Delaney, Mellon Bank Strategic Technology Group.

The first paper dealt with a marketing problem: matching possibly new customers with newly released musical CDs. The large but statistically sparse data set available for "training" was a history of individual purchases and rejections for many customers. The novel methodology introduced statistical techniques for creating a simultaneous clustering of both customers and products in order to forecast customer interest in a new product. The statistical tools employed iterative methods for improving the clustering in one dimension, e.g., customer type, based on the latest assessments in the other dimension, e.g., record type, with alternation between the two dimensions.

The second presentation used a combination of cluster analysis and query matching in order to assess the adequacy of a reduction from a relational to a flat database. Specifically, a reduced (that is, flat) database was subjected to a cluster analysis. These clusters were then matched by (optimal) queries made on the original (that is, relational) database. The better the matching, the better the reduction. The novel methodology employed neither statistical nor iterative procedures.

The third presentation reviewed a bank's case-study of how to build clustering models of customers in order to forecast attrition of credit-card use. The authors reviewed tradeoffs that they faced in choosing among different datamining tools. They considered, for example, the relative costs in terms of speed of computation, price for extracting warehoused data, etc., versus the accuracy of the clustering tool.

These presentations, as well as the lively exchanges that they provoked, suggest the following [two] high-level areas for future research:

1. What is the meta-analysis appropriate for combining the outputs from different datamining clustering tools, especially in cases where the tools use incomparable methodologies, e.g., as in a contrast between the statistically based tools of paper 1 and the databased tools of paper 2? That is, rather than trying to choose among competing datamining tools for cluster analysis, how may such tools be used simultaneously?
2. What is the theory of (optimal) iterative methods to be used with clustering algorithms?

Workshop 5: Visual Methods for the Study of Massive Data Sets

Organizers: Bill Eddy and Steve Eick

Scribe: Ashish Sanil

This is a very brief outline of some of the points that were raised in the CONALD Workshop on Visual Methods for the Study of Massive Data Sets

Visual Methods as an Aid for Modeling Data

Visual methods for understanding data seem to be most useful when the goals of the analysis are not well-specified or when no precisely quantifiable models may be evident. Visual methods help one to:

1. Formulate precise models for the data.
2. Examine the output produced by the models.
3. Perform diagnostic tests to validate the models. (The diagnostics might suggest more appropriate/refined models.)

Multiple Ways of Displaying the Data

Most of the talks emphasized the need to be able to examine different aspects of the data through a wide variety of graphical displays. A good system for providing multiple graphical views of the data should have the following properties:

1. The views should complement each other to reveal different aspects of the data.
2. The views should be linked to each other.
3. Each display should have highlighting/subset-selection features (with the highlighted points or selected subset being highlighted or selected in all other other linked views also).
4. Views should be available at various levels of aggregation from individual records to the entire dataset.
5. The system should be able to present displays of datasets which have a more complicated (hierarchical) structures than the usual flat-file datasets.

Designing Better Displays

There was a good bit of discussion on designing displays which would convey information efficiently through graphics which are intuitive easy to understand. Suggestions for further work along these lines were:

1. Consult with cognitive scientists, psychologists, etc. to understand efficient use of color, plotting symbols, layouts, dynamic displays, and so on.
2. Focus on the end-user. Explicitly consider the scientific area, organizational culture, computing environment, etc. to design more appropriate and easily understandable displays.

Interplay Between Visualization, Modeling, and Mining

There is an important interplay between data mining and modeling techniques aimed at discovering patterns and visual techniques aimed at understanding. Visualization techniques are inherently less scalable than modeling algorithms since they are bounded by limits in human perception. Combining techniques can be complementary:

1. Exploratory visualization may guide model selection.
2. Statistical algorithms such as smoothing help may visual displays more understandable.

Intelligent Features

We need to create visualization systems with intelligent features to assist the naive user. Some specific suggestions were:

1. Conduct systematic experiments to understand how experts conduct graphical data analysis. Try to identify principles of visual data analysis and incorporate them in a system which could guide the user and suggest analyses/views based on these principles.
2. Construct a taxonomy of datasets by their structural properties. Then devise generic strategies for the visual display of each kind of dataset.
3. Categorize datasets by discipline: financial data, engineering data, etc. Then consult with experts in the respective disciplines to identify the kind of visual analyses most appropriate for that discipline.

Massive Datasets

Discussion on problems particular to massive datasets led to a classification of the problems into two broad categories:

1. *Hardware Limitations*: We need some hardware solutions to the kind of problems one would encounter when we try to display a larger number of quantities than the screen resolution will permit.
2. *Limits of Human Perception*: We need to address the problem of how to display larger quantities of data than what is possible for people to discern and process.

Visual Methods for Data Analysis Vs. Data Presentation

Several participants pointed out the the need to distinguish between using graphical methods for exploring and understanding data with graphical methods for highlighting and presenting certain features of the data. It might be beneficial to keep keep this distinction in mind while designing new systems.

Promising Research Topics

Some promising research topics discussed include:

1. Visual techniques for real-time data collection systems.
2. Techniques to combine visualization and modeling.
3. The creation of software environments for building interactive visualizations.
4. Techniques for scaling visualization to the size datasets that are readily manipulated on inexpensive PCs.

Workshop 6: Learning from Text and the Web

Organizers: Yiming Yang, Jaime Carbonell, Steve Fienberg and Tom Mitchell

Scribe: Mark Craven

An increasing fraction of the world's information and data is now represented in textual form. For example, the World Wide Web, online news feeds, and other Internet sources contain a tremendous volume of textual information. The goal of the CONALD workshop on *Learning from Text and the Web* was to explore computer methods for automatically extracting, clustering and classifying information from text and hypertext sources.

The workshop included ten oral paper presentations, an organized discussion by a panel of distinguished researchers, and a handful of other contributed papers. The workshop provided a good survey of the state of the art in machine learning methods applied to text processing tasks. The presented work involved a wide array of learning approaches, including finite-state-machine induction [8, 18], neural networks that can accept *advice* from users [22], relational learning methods [17, 23], statistical clustering algorithms [6, 7, 13, 24], boosting methods [1], algorithms for learning with hierarchical classes [7, 16], and active learning methods [14, 19]. A principal limitation of many of these approaches is that they do not directly reflect attempts to develop formal models of the text phenomenon of interest.

The research presented at the workshop also spanned a broad range of application tasks, including: information extraction [9, 8, 12, 17, 18], information finding [22], information integration from Web sources [18], automatic citation indexing [2, 10], event detection in text streams [24], document routing [1] and classification [5, 17], organization and presentation of documents in information retrieval systems [6, 7], collaborative filtering [3], lexicon learning [4], query reformulation [11], text generation [21] and analysis of the statistical properties of text [15]. In short, the state of the art in learning from text and the web is that a broad range of methods are currently being applied to many important and interesting tasks. There remain numerous open research questions, however.

Broadly, the goals of the work presented at the workshop fall into two overlapping categories: (i) making textual information available in a structured format so that it can be used for complex queries and problem solving, and (ii) assisting users in finding, organizing and managing information represented in text sources. As an example of research aimed at the former goal, Muslea, Minton and Knoblock [18] have developed an approach to learning *wrappers* for semi-structured Web sources, such as restaurant directories. Their method is able to induce extraction rules from small numbers of labeled examples. These learned extraction rules are then applied so that Web pages can be treated like structured databases. As an example of work geared toward the latter goal, Shavlik and Eliassi-Rad [22] have developed an approach to increasing the communication bandwidth between users and learning agents that perform tasks such as home-page finding. Their approach enables a user to give advice to a learning agent at any time during the agent's lifetime. The advice is incorporated into the agent's learned model where it may be subsequently refined by reinforcement-learning methods.

The likely future impact of research in text learning is twofold. First, much of the information that is currently available only in text form will automatically be mapped into a structured format. This ability would mean that the Web queries would not be limited to keyword searches for individual documents relevant to the query. Instead, we could directly get answers to queries whose answers are distributed across multiple Web sources. Moreover, this ability would mean that text sources, such as the Web, could be used for planning and problem solving (e.g. an agent that would use information on the Web to make travel plans for IJCAI-99). The second likely impact of continued work on

learning from text and the Web is greatly improved methods for finding, organizing, and presenting information in free-text data sources including Web pages, emails, transcribed radio/TV broadcasts and newswire stories.

The workshop brought to light numerous promising research topics and raised many open research questions for further exploration:

- *Learning from structure in hypertext.* There are many opportunities to go beyond bag-of-words representations documents by exploiting HTML formatting and patterns of connectivity in hypertext.
- *Developing statistical models that represent hypertext structure.* In addition to developing algorithms that are able to exploit document structure, it is also important to develop well-founded statistical models for such tasks. For example, we might develop stochastic models for graphs of linked objects by exploiting ideas from the literature on *social networks*.
- *Exploiting NLP techniques more in text learning.* Much of the current work in text learning concentrates on word occurrence statistics and ignores other linguistic structure. There is much work to be done in understanding how natural language processing methods can be applied to gain more accurate learners.
- *Learning from temporal patterns in text.* Some text processing tasks, such as topic tracking and event detection, involve a stream of text data over time. New methods are needed to represent, detect and exploit content (and Web structure) that changes over time.
- *Learning from available domain knowledge, advice and data.* In addition to learning from training data, text-learning systems should be able to take advantage of background knowledge and on-line advice offered by users.
- *Systems that learn to improve the organization and presentation of information.* As an example of this type of system, Perkowitz and Etzioni [20] have begun developing adaptive Web sites. Other related issues include learning to recognize certain types of queries, and learning from passive relevance information.
- *Developing statistical models for interaction data.* In addition to developing effective methods for systems such as adaptive Web sites, it is also important to develop statistical models of non-stationary user interaction data.
- *Learning text classifiers when word statistics are sparse.* Current research in this area is exploring such methods as term clustering, feature reduction, document clustering, active learning, and using hierarchical class information.
- *Combining evidence from multiple sources.* This issue is relevant to information extraction, information retrieval, and text classification. In information retrieval, for example, we may want to combine document rankings that consider such factors as document-query similarity, document popularity, and editorial vetting of documents.

Current research in learning from text and the Web is already quite cross-disciplinary. As the above list of promising research topics indicates, however, the hard problems in the area call out for expertise in a wide variety of disciplines including machine learning, statistics, information retrieval, natural language processing, planning, and human-computer interaction.

A challenging question in the inter-disciplinary research is, can we significantly improve the state of the art by introducing more principled formal models about text phenomena of interest? Furthermore, how can we determine the suitability of such models? Empirical evaluation (e.g. cross-validation) **alone** may not be sufficient for analyzing the behavior and limitations of algorithms. On the other hand, comparing models without evaluating their practical impact in specific tasks is clearly inadequate. This workshop presented a promising trend of research that addresses both theoretical and empirical concerns: Hofmann and Lafferty and Venable both proposed new statistical language models for document clustering. Apte presented strong empirical evidence for using boosting to improve the state of the art in text categorization. These exciting research findings encourage further investigation for more satisfactory answers in the future.

References

- [1] C. Apte, F. Damerau, and S. Weiss. Text mining with decision rules and decision trees.
- [2] K. D. Bollacker, S. Lawrence, and C. L. Giles. CiteSeer: An autonomous system for processing and organizing scientific literature on the Web.
- [3] O. de Vel and S. Nesbitt. A collaborative filtering agent system for dynamic virtual communities on the Web.
- [4] A. F. Gelbukh, I. A. Bolshakov, and S. N. Galicia-Haro. Automatic learning of a syntactical government patterns dictionary from Web-retrieved texts.
- [5] B. Gelfand, M. Wulfekuhler, and W. F. Punch III. Automated concept extraction from plain text.
- [6] M. Goldszmidt and M. Sahami. A probabilistic approach to full-text document clustering.
- [7] T. Hofmann. Learning and representing topic.
- [8] C.-N. Hsu and M.-T. Dung. Wrapping semistructured Web pages with finite-state transducers.
- [9] Hull and Fluder. Text mining the Web: Extracting chemical compound names.
- [10] A. Kehagias and V. Petridis. Automated building of a database of neural network papers.
- [11] Y. S. Kwon and N. H. Kim. The effect of relevant input information in ID3's learning performance.
- [12] Z. Lacroix and A. Sahuguet. Information extraction and heuristics for human-like browsing.
- [13] J. Lafferty and P. Venable. Simultaneous word and document clustering.
- [14] R. Liere and P. Tadepalli. Active learning with committees: Preliminary results in comparing winnow and perceptron in text categorization.
- [15] P. P. Makagonov and M. A. Alexandrov. Tool for measurement of statistical properties of text.
- [16] D. Mladenic and M. Grobelnik. Feature selection for classification based on text hierarchy.
- [17] R. J. Mooney. Learning for information extraction, querying, and recommending.

-
- [18] I. Muslea, S. Minton, and C. Knoblock. Wrapper induction for semistructured Web-based information sources.
 - [19] K. Nigam and A. McCallum. Pool-based active learning for text classification.
 - [20] M. Perkowitz and O. Etzioni. Adaptive web sites: an ai challenge. In *Proceedings of the Fifteenth International Joint Conference on Artificial Intelligence*, Nagoya, Japan, 1997. Morgan Kaufmann.
 - [21] D. R. Radev. Learning correlations between linguistic indicators and semantic constraints: Reuse of context-dependent descriptions of entities.
 - [22] J. Shavlik and T. Eliassi-Rad. Building intelligent agents for Web-based tasks: A theory-refinement approach.
 - [23] S. Slattery and M. Craven. Learning to exploit document relationships and structure: The case for relational learning on the Web.
 - [24] Y. Yang, T. Pierce, and J. Carbonell. Event detection.

Workshop 7: Robot Exploration and Learning

Organizers: Howie Choset, Maja Matarić and Sebastian Thrun

Scribes: Nicholas Roy and Wes Huang

1 Introduction

Recently, robot exploration and learning techniques have begun to have a profound impact on the field of robotics. Exploration and learning techniques have been at the heart of many recent successes in the field of robotics, where they have provided additional robustness to robot behavior as well as eased the task of robot programming.

The workshop “Robot Exploration and Learning” brought together leading researchers to discuss open questions in robot exploration and learning, including: How can we scale up robot exploration and learning so that more can be learned from less data? How can we devise efficient exploration strategies for dynamic and high-dimensional environments? What role will learning and exploration play in upcoming application domains such as service robotics? How can learning support cooperation of multi-robot teams? How can we bring together and achieve synergy between the various disciplines involved in the study of robot exploration and learning and beyond?

This summary is organized in three parts: a chronological summary of workshop events, and a thematical summary that responds to the central issues addressed by the CONALD meeting, and a keyword list of items that represent interesting open research topics in the field, compiled in the workshop.

2 Minutes

Before the talks began, we collected some questions to be answered during the day

1. Why is robot learning and exploration so hard?
2. What will our robots do?
3. Why should they do it?
4. How shall our algorithms we represent information?
5. What are the most promising learning methods?

Leslie Kaelbling gave the first invited talk. Her theme was that learning algorithms do not work for autonomous robots because:

- Learning is off-line
- Supervised learning is a batch-job
- Reinforcement learning runs in simulation, but in the real world it is too slow

She asserted that we need true on-line learning and then reviewed some methodologies for performing it. Three options were given: neural nets, nearest neighbor, and reinforcement learning:

- Neural nets rarely work successfully in online search because they have trouble with non-stationary sampling distributions. For example, a robot would have to forget what it learned in a hallway in order to adapt for a large open room.

- Nearest neighbor solutions are good because if the original query were wrong, you know the answer is in a neighborhood of the correct answer. Also, nearest neighbor solutions can adapt to different localities, i.e., changes in distribution. However, nearest neighbor has trouble with dynamic and highly dimensioned environments.
- Reinforcement learning is model-free but uses experiences differently. Also, it explores unsystematically and therefore requires supervision

Kaelbling suggested that we sacrifice optimality for computational efficiency, learning efficiency, and representational efficiency. Also to use variable stochastic spatio-temporal resolutions for problem representation.

Peter Stone then described his PhD work on the multi-agent soccer problem. Current reinforcement learning techniques do not provide an appropriate solution to his multi-agent problem because the teammates' actions are largely hidden from one another and their policies can change. Christoph Zack stated that reinforcement learning does work in a multi-agent setting, except under some very strong assumptions (e.g., only one learner). Claude Touzet talked about cooperative robotics as well. He was interested in implementing communication policies among the robots and how to distribute a global reward to all the robots when their performed well. He also considered how to increase the duration of cooperative robotics experiences. Maja Mataric pointed out that some progress has already been made in using communication to alleviate the hidden state and credit assignment problems in multi-robot learning.

Pallela described his groups work on three-dimensional viewing using an entropy loss model. Michael Littman used POMDP's for robot exploration of an unknown environment. He directed a Khepera robot around a simple walled environment; the POMDPs then inferred the structure of the outer walls by looking for the appropriate sequence of left and right turns. The final talk was delivered by Daniel Nikovsky, who emphasized the importance of structured probabilistic reasoning in robot perception.

Maja Mataric gave the second invited talk, on Biological Inspirations for Facilitating Reinforcement Learning in Challenging Domains. There are three reasons for investigating this problem:

- insight into evolution and natural behavior
- intuition as to how lessons in evolution can help us build better robots, and
- learning how to build robots that better mimic life-like behaviors

Mataric suggested that there is a great expectation with current learning algorithms: so little is put in and we expect so much to come out. From our poor sensors, poor domain knowledge, poor robot morphology, and poor feedback, we want optimality, provability, real-time execution, and life-time learning. Instead of worrying about optimality, we should consider efficiency; likewise, instead of worrying about reductionist simplicity, we should consider realistic complexity instead.

She mentioned some examples of applied biological inspiration: topological mapping, motor primitives schemas, shaping in reinforcement learning (i.e., providing intermediate feedback), use of communication, and imitation, all of which can be effectively applied to obtain more robust and efficient robot learning.

Hideki Asoh spoke about the natural interaction between people and robots in real-world environments, such as the office and home. He distinguished between explicit and implicit interaction. Examples of explicit interaction are dialogue, showing, doing where as implicit interaction includes gazing pointing, etc. Andrew Baynell also talked about reinforcement learning with robots and showed a demonstration of three small robots, equipped only with proximity sensors, performing obstacle avoidance. Malcolm Ryan spoke about decomposition in reinforcement learning (RL), in which RL is not used to combine behaviors, but rather to build them. Jefferson Coelho discussed

control laws.

Noel Sharkey asked why we don't copy nature. He asserted that it is difficult to copy nature and should focus on the incremental development of controllers. Finally, Ben Krose described a method that models the environment in a sensor configuration space using Kohonen networks.

The end of the session was spent in open discussion. The following are some of the questions that were asked:

1. how do we insert domain knowledge and geometry into reinforcement learning problems?
All agreed perception was a big problem and questioned how much (percentage-wise) of our actions are learned. The hack-it-together and see-what-happens method has produced some results but perhaps we should be looking at problems more intelligently.
2. how do you do optimization in real time?
3. how do you tell your robot is learning?
4. what are we trying to do?
5. how hard is it to evaluate in complicated domains?

Tony Stentz gave the final invited talk on some of the robotics projects at CMU. Some key robotic tasks he suggested are:

- Vehicle scheduling
- Material flow
- Search & rescue
- Sweeping & clearing
- Coverage
- Exploration
- Guarding

There are many types of performance criteria for robotics; some include:

- Optimality
- Completeness
- Real-time
- Uncertainty
- Reactive (dynamic)

He talked about his approach to dynamic programming and incremental repair of plans. It was asked during his talk if navigation was a solved problem. For the applications that Tony worked on, the answer is yes, but in general, no.

William Smart spoke about reinforcement learning on a real robot and described many problems associated with the approach: huge state space, hidden state, little data, exploration risks (such as falling down stairs), and realistic time-frames. Bruce Digney then talked about simple hierarchy learning.

Brian Yamauchi talked about frontier-based exploration and continuous localization while exploring, with plans to extend this approach to a topological and metric representation. Howie Choset then talked about a topological representation, termed the Generalized Voronoi Graph (GVG), which captures the salient geometry of the environments and is used to direct the robot to explore new regions and localize itself, at distinct times. Jim Jennings described his use of an augmented GVG to perform distributed map making in dynamic environments.

Finally, Sebastian Thrun described a spectrum ranging from programming to learning, where machine learning was at one end and conventional programming was at the other. He argued for

the importance of integrating both: While the predominant methodology for making robots work, an inability to adapt to their environments imposes serious scaling limitations. Learning, on the other hand, is promising, but must existing learning methods learn tabula rasa. In domains such as robotics, it is infeasible to learn complex robot controllers from scratch. Thrun argued that new research is needed, on methods that are flexible enough to integrate both, to combine the strengths of conventional programming with that of machine learning methods.

3 Answers to CONALD Questions

1. What is the state of the art? Examples of successful systems? Best answer seems to be indoor navigation. There are almost unanimous consensus that indoor office navigation is a solved problem.
2. Long-term goals? This issue was not discussed in depth. Ben Krose's dream of showing robot coffee pot picture, and having robot find coffee pot, might qualify.
3. Promising research topics: see section on thesis topics. Kaelbling and Mataric both emphasized the need to worry less about optimality and formalisms, and allow synthesis to drive analysis. Possible need for paradigm shift? Tony Stentz offered many very applied research topics, in such areas as material flow, coverage, search and rescue, etc. One topic that appeared often was the importance of representations - possible research topic in formal analysis of the optimality of different spatio-temporal representations.
4. Opportunities for cross-disciplinary research
Sharkey suggested robots need better, more realistic bodies. Mataric suggests that biology has a lot to offer. A possible that the reason biology has not had more of an impact is that we have not abstracted the right principles from biology, and we have not given the systems decent actuators and sufficient complexity.

4 Issues that might lead to Thesis Topics

This section sketches open aspects that could form promising problems for Ph.D. theses on robot exploration and learning.

1. Really Online Systems
 - (a) efficiency
 - (b) one-shot
 - (c) noise tolerant
 - (d) insensitive to non-stationary sampling data
2. Automatic Development of Hierarchical Decomposition
 - (a) spatial
 - (b) general state space
 - (c) motor actions
3. Living one (long) life
 - (a) changing body
 - (b) changing environment
 - (c) changing goals

4. Partial Observability
 - (a) learning what to remember
 - (b) active information gathering
 - i. learning where to move to acquire more info
 - ii. learning where to move to refine or improve current map (e.g. localization)
5. Multiple-robots
 - (a) learn to cooperate
 - (b) learn to compensate
 - (c) learn to communicate
 - (d) learn to achieve group goals
 - (e) scaling up to very large environments
6. Building in prior information
 - (a) world dynamics
 - (b) utility
 - (c) behavior
7. How to represent environment
 - (a) probabilistic
 - (b) relational
 - (c) geometric
 - (d) algebraic
8. Perception
 - (a) tune low-level mechanism intelligence
 - (b) attention
 - (c) fusion
 - (d) models (sensor models)
9. Satisficing
 - (a) formalizing and striving for efficiency over optimality
 - (b) real-time well-defined satisficing solutions

While this list is sketchy and only hints at specific problems, there was an agreement among the participants that these are some of the key problems of our current best methods. Research along any of those lines is likely to produce new and interesting insights into scaling up robot exploration and learning.