

Blackboard-based Sensor Interpretation using a Symbolic Model of the Sources of Uncertainty in Abductive Inferences

Norman Carver and Victor Lesser

Department of Computer and Information Science
University of Massachusetts, Amherst, Massachusetts 01003
(carver@cs.umass.edu, lesser@cs.umass.edu)

Abstract

Sensor interpretation involves the determination of high-level explanations of sensor data. The interpretation process is based on the use of abduction. Interpretation systems incrementally construct hypotheses using abductive inferences to identify possible explanations for the data and, conversely, possible support for the hypotheses. We have developed and implemented a new blackboard-based interpretation framework called RESUN. One of the key features of RESUN is that it uses a model of the sources of uncertainty in abductive interpretation inferences to create explicit, symbolic representations (called SOUs) of the reasons why hypotheses are uncertain. The symbolic SOUs make it possible for the system to understand the reasons why its hypotheses are uncertain so that it can dynamically select the most appropriate methods for resolving uncertainty. Our model of uncertainty defines a set of classes of SOUs that are applicable to interpretation problems which can be posed as abduction problems. Each interpretation application may require slightly different instances of each of the classes of SOUs to best represent uncertainty. We have implemented the RESUN framework using a simulated aircraft monitoring system and have run experiments that demonstrate how the SOUs enable the use of more effective interpretation strategies. To verify the generality of the approach, we are also using RESUN to implement a sound understanding testbed.

Introduction

Sensor interpretation involves the determination of high-level *explanations* of sensor and other observational data. The interpretation process is based on a hierarchy of *abstraction types* like the one in Figure 1 for a vehicle monitoring application. An interpretation system incrementally constructs *hypotheses* that represent possible explanations for subsets of the data. For example, in vehicle moni-

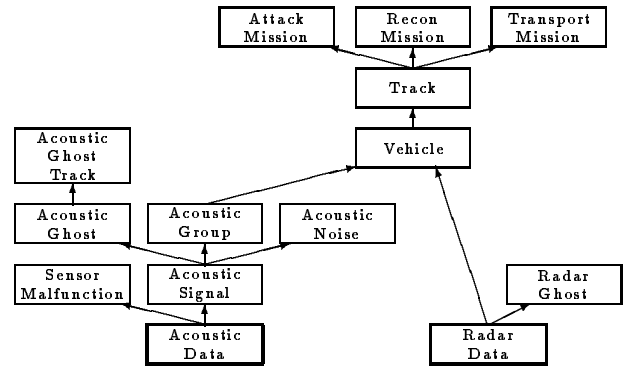


Figure 1: Vehicle monitoring abstraction hierarchy.

toring, data from sensors (e.g., Acoustic Data and Radar Data) is abstracted and correlated to identify potential vehicle “sightings” (Vehicle hypotheses), vehicle movements (Track hypotheses), and vehicle goals (Mission hypotheses).

The basis for interpretation is the concept of *abduction*: if B 's can *cause* A 's then given some A , a_i , we might hypothesize that there is some B , b_j , that is an *explanation* for a_i . Conversely, given a possible B , b_k , the existence of A 's that b_k could have caused provide *support* for b_k . An interpretation type hierarchy like the one in Figure 1 is effectively a *causal hierarchy*. For example, the existence of an Attack-Mission causes there to be a vehicle Track with certain characteristics, the Track causes there to be a sequence of Vehicle sightings with appropriate parameters, and this eventually causes there to be appropriate signals that result in sensor data. Thus an interpretation system makes *abductive inferences* that identify possible explanations for the data and possible support for its hypotheses.

We have developed and implemented a new blackboard-based framework for sensor interpretation called RESUN. RESUN uses a model of the uncertainty in abductive interpretation inferences to create explicit, symbolic representations of the *sources of uncertainty* in hypotheses. For example, a Track hypothesis in an aircraft monitoring system may be uncertain because its supporting sensor data

*This work was supported by the Office of Naval Research under University Research Initiative grant number N00014-86-K-0764.

might have alternative explanations as a ghost or as part of a different aircraft Track or it may be uncertain because its evidence is incomplete or because it is not known whether there is a valid Mission-level explanation for the Track; these are some of the possible sources of uncertainty for Track hypotheses. As interpretation inferences are made in RESUN, symbolic statements, *SOU*s, are attached to the hypotheses to represent their current sources of uncertainty.

Interpretation problems have often been approached using *blackboard* frameworks. However, most existing blackboard-based interpretation systems have been limited to *incremental hypothesize and test* strategies for resolving uncertainty in hypotheses instead of being able to use more powerful *differential diagnosis* strategies [4]. The RESUN framework is capable of supporting a wide variety of interpretation strategies—including differential diagnosis strategies. Because the symbolic *SOU*s allow the system to understand the *reasons* why the hypotheses are uncertain, they make it possible for the system to dynamically select methods for resolving its uncertainty instead of being limited to a fixed strategy like hypothesize and test. In RESUN, interpretation is viewed as an incremental process of gathering evidence to resolve particular sources of uncertainty in the hypotheses. RESUN's planning-based control mechanism is described in more detail in [3, 4].

Hypotheses and Extensions

Before we examine the representation of uncertainty, we must introduce another aspect of the framework. In RESUN, a hypothesis is viewed as a set of *extensions*, each representing a different possible version (or set of versions) of the hypothesis. Because interpretation requires constructive problem solving (see the section on interpretation vs. classification below), the hypothesis versions of interest are identified as a part of the problem solving process. As evidence is gathered for a hypothesis, the values of the hypothesis' parameters are defined by the parameters of the data and hypotheses that make up its support and explanation evidence. For example, Vehicle (sighting) hypotheses not only can support Track hypotheses, they also constrain the values of the aircraft ID and positions parameters of a Track they support.

However, for any set of evidence, the values of the hypothesis' parameters may be uncertain—i.e., the evidence may only partially constrain the parameters. We handle parameter uncertainty by allowing the use of sets or ranges to represent the potentially correct values for a parameter and by allowing array-valued parameters to be incompletely specified (e.g., aircraft positions over time). For example, the value of the ID parameter of a Track hypothesis may be represented as a set of possible

values and the positions parameter may be incomplete (representing complete uncertainty about the positions of the aircraft at certain times). Because parameter values may be uncertain, every time evidence is added to a hypothesis it may further constrain the possible parameter values (see the discussion of Figure 2 below). In other words, gathering evidence for an interpretation hypothesis not only *justifies* the hypothesis, it may also *refine* it by further constraining its parameter values.

Because most interpretation evidence is uncertain and because there can be multiple instances of any data or hypothesis types, it is possible for there to be alternative pieces of evidence for a hypothesis. For example, there might be two different Vehicle hypotheses for time t_i that are both consistent with an existing (partial) Track hypothesis. Because alternative evidence may refine a hypothesis differently and because each alternative is uncertain, multiple versions of the hypothesis may have to be maintained. When these versions are maintained as independent hypotheses—as they are in most blackboard-based interpretation systems—valuable information about the relationships between the versions is lost. Instead, we maintain alternative hypothesis versions as different extensions of a single root hypothesis. This allows us to understand, for example, that a Track *hypothesis* may very likely be correct even though we are still uncertain about the correct *extension* of the hypothesis—i.e., we are quite certain that there is an aircraft in the monitored region, but we are uncertain about its exact path.

Evidence and Uncertainty

As we have already stated, the basis of the interpretation process is *abduction*. An interpretation system makes abductive inferences that identify possible explanations for data and, conversely, possible support for hypotheses. Abductive inferences are uncertain rather than logically correct inferences. In other words, abductive interpretation inferences provide *evidence* for the hypotheses rather than conclusively proving them.

For each interpretation type T , the interpretation specification defines the type's support, S_T , and its possible explanations, E_T . The support, S_T , is a set of support *sources*—i.e., $S_T = \{S_k\}$. Each support source consists of a set of type instance specifications—i.e., for each $S_l \in S_T$, $S_l = \{S_{il}\}$ where each $S_{jl} \in S_l$ is a type *instance* specification. By type instance specification, we mean an interpretation type along with parameter constraints—e.g., a Vehicle type with constraints on the position and ID parameters. For example, a Track hypothesis constrains the aircraft IDs of its supporting Vehicle hypotheses to be identical and their positions to be consistent with the movement characteristics of the particular aircraft. This definition of the sup-

port, $\mathbf{S}_T$, as a set of support sources, $\{\mathbf{S}_k\}$, is done to model domains like aircraft monitoring where there may be multiple *sources of evidence* for some types—e.g., a Vehicle hypothesis may be supported by radar data or by a set of Group hypotheses based on acoustic sensor data (see Figure 1).

The possible explanations, $\mathbf{E}_T$, is a set of types, $\{E_i\}$, each of which might explain some type T hypothesis. For example, a Track hypothesis might be able to be explained as an Attack-Mission, a Recon-Mission, or a Transport-Mission; these are the three possible explanation types for the Track type. Note, though, that because of the constraints which each of these Mission types places on the vehicle ID and positions of associated Tracks, each particular Track hypothesis may only be able to be explained by some subset of these Missions types. Based on these definitions, every hypothesis of type T , H_T , is the result of a set of abductive inferences each of which is of the form: $\{H_{S_{j_l}} \Rightarrow H_T\}$ (support evidence) or $H_T \Rightarrow H_{E_j}$ (explanation evidence) where $H_{S_{j_l}}$ is a hypothesis corresponding to type instance specification $S_{j_l} \in \mathbf{S}_l$, $\mathbf{S}_l \in \mathbf{S}_T$ and H_{E_j} is a hypothesis of type $E_j \in \mathbf{E}_T$.

Our symbolic representation of interpretation uncertainty is based on a model of the underlying uncertainties in abductive inferences and on the requirements for controlling interpretation systems—i.e., that the system be able to identify the methods it could use to resolve its uncertainty. The basic source of interpretation (abduction) uncertainty is the possibility of alternative explanations for data; uncertain data and incomplete models of the possible explanations prevent the direct, conclusive determination of the interpretations of the data. However, there are factors other than the possibility of alternative explanations for data that influence the level of belief in hypotheses. As a result, there are several ways for interpretation systems to go about resolving uncertainty. In order to enable the use of all possible methods, our symbolic SOUs represent more information than just the possible alternative explanations for data.

Hypothesis correctness can only be guaranteed by discounting all of the possible explanations for the supporting data—i.e., doing complete differential diagnosis; even if complete supporting evidence can be found for a hypothesis there may exist alternative explanations for all of this support. However, while complete support cannot guarantee hypothesis correctness, the amount of supporting evidence is often a significant factor when evaluating the belief in a hypothesis (this is the basis of hypothesize and test strategies). For example, once an aircraft Track hypothesis is supported by sensor data from a “significant” number of (correlated) individual sensor sightings, the belief in the track may be fairly high regardless of whether alternative explanations for its supporting data are still possible. In addition,

complete differential diagnosis is typically very difficult because it requires the enumeration of all of the possible interpretations which might include the supporting data—many of which may not be able to be conclusively discounted. Thus, a combination of hypothesize and test and (partial) discounting of critical alternative explanations must be used to gather sufficient evidence for interpretation hypotheses; our SOU representation is designed to drive this sort of process.

The model of uncertainty in interpretation inferences is based on the basic uncertainty of abductive inferences, the factors which affect the belief in hypotheses, the alternative methods for resolving uncertainty, and our extensions model of hypotheses. The model specifies a set of potential *classes* of SOUs for any hypothesis extension (see [3] for more complete definitions). Particular instances of these SOU classes will be instantiated for each application and for each set of inferences supporting an interpretation hypothesis. We have defined the following potential *classes* of SOUs for a hypothesis extension H_T :

- **partial evidence** – Denotes the fact that there is incomplete evidence for the hypothesis. For example, a No Explanation SOU means that no explanation has been determined and a Partial Support SOU means that for some support source, l , the current set of support hypotheses, $\{H_{S_{j_l}}\}$ is incomplete—i.e., $\{S_{j_l}\} \subset \mathbf{S}_l$ and $\{S_{j_l}\} \neq \mathbf{S}_l$. For example, a Track hypothesis which has not yet have been examined for valid mission-level explanations will have an No Explanation SOU associated with it. typically have incomplete supporting Vehicle hypothesis evidence (no supporting Vehicle hypotheses for some times included within Track).
- **possible alternative support** – Denotes the possibility that there may be alternative evidence which could play the same role as a current piece of support evidence—i.e., that there exists a hypothesis $H'_{S_{j_l}}$ which is the correct S_{j_l} support rather than $H_{S_{j_l}}$. This reflects the fact that though “a hypothesis” may be quite certain, there can still be uncertainty over the correctness of individual pieces of evidence for the hypothesis—i.e., uncertainty over the correct extension. This is an additional complication for differential diagnosis in interpretation problems as compared with classification problems (see Section 2). Classification problems do not have to contend with multiple instances of the types and so do not have this source of uncertainty. Interpretation must consider the possibility that there is alternative supporting evidence for a hypothesis—i.e., that there is a different *version* of the hypothesis which is actually correct.
- **possible alternative explanation** – Denotes the possibility that there may be alternative ex-

planations for the hypothesis—i.e., that there exists a hypothesis H'_{E_k} , $E_k \in \mathbf{E}_T$ which is the correct explanation rather than H_{E_j} . These SOUs explicitly identify the possible explanation types based on the characteristics of the hypothesis.

- **alternative extension** – Denotes the existence of a competing, alternative extension of the same hypothesis. In other words, an alternative version of the hypothesis has been created using one or more pieces of evidence that are inconsistent with the existing versions of the hypothesis—i.e., using alternative support and/or an alternative explanation. This is the primary representation of the relationships between hypotheses.
- **negative evidence** – Denotes the failure to be able to produce some particular support evidence, S_{ji} , or to find any valid explanations in $\mathbf{E}_T$. Negative evidence is not conclusive because it also has sources of uncertainty associated with it—e.g., that sensors may have missed some data.
- **uncertain constraint** – Denotes that a constraint associated with the inference could not be validated because of incomplete evidence or uncertain parameter values. This SOU represents uncertainty over the *validity* of an evidential inference whereas the other SOUs are concerned with the *correctness* of inferences. See the example described below for further explanation of this SOU.
- **uncertain evidence** – Technically, this is not another source of uncertainty *class*. Uncertain evidence SOUs merely serve as placeholders for the uncertainty in the evidence for a hypothesis because the sources of uncertainty are not automatically propagated as evidential inferences are made. They denote the fact that an evidential inference is uncertain because the inference contains uncertain constraint SOUs and/or the hypothesis extension which is the basis for the inference contains SOUs.

Figure 2 shows three extensions of a Track hypothesis along with their associated SOUs and parameters. Track-Ext₁ is an *intermediate extension* while Track-Ext₂ and Track-Ext₃ are *alternative maximal extensions*. The alternative extensions result from competing possible explanations of the Track as an Attack-Mission or as a Recon-Mission. This alternatives relationship between these Mission hypotheses is represented by the *alternative extension* SOUs in Track-Ext₂ and Track-Ext₃. These SOUs indicate that there is a negative evidential relationship between the extensions: more belief in Track-Ext₂ or Attack-Mission results in less belief in Track-Ext₃ or Recon-Mission (and vice versa). They also make it possible for the system to recognize that the uncertainty in Attack-Mission need not be directly resolved, but can be pursued by resolving the uncertainty in Recon-Mission or by resolving the uncertainty in the Track’s parameter values (in order to limit its consistent interpreta-

tions). This example also demonstrates how extensions represent different versions of hypotheses: the uncertainty in the value of Track-Ext₁’s ID parameter has been resolved differently by the alternative explanations. The uncertainty that results from each explanation only being consistent with a subset of the possible values for the Track’s ID parameter is represented by *uncertain constraint* SOUs. These SOUs do not appear in the figure because they are maintained as part of the inferences; they are accessed through the “placeholder” *uncertain explanation* SOUs which represent the overall uncertainty in the explanations.

Numeric Summarization of SOUs

In addition to the symbolic uncertainty encoding, RESUN’s evidential representation system also includes a framework for numerically summarizing the symbolic SOUs in the evidence for a hypothesis. The summarization process produces a composite characterization of the uncertainty in a hypothesis in terms of an overall belief rating and the relative uncertainty contributions of the different classes of SOUs (listed above). This composite numeric summary is used in evaluating the satisfaction of termination criteria and in selecting the hypotheses to pursue and the methods to use to pursue them. Having a composite rating allows more detailed reasoning about termination and focusing decisions than would be possible with a single number rating. For example, it can distinguish between a hypothesis that has low belief due to a lack of evidence having been gathered for it and one for which there is negative evidence—i.e., evidence that it is incorrect. It can also show whether residual uncertainty results from actual competing hypotheses (that may need to be examined further) or whether it is simply due to the *possibility* of alternative explanations, etc.

The summarization process evaluates the SOUs for a hypothesis extension using evaluation functions specific to each interpretation type. Hypothesis extensions are summarized by rating the relative contribution of each SOU to the uncertainty of the extension and then using a combining function to produce the composite rating. The placeholder *uncertain evidence* SOUs are evaluated by evaluating the evidence they represent. This results in a recursive summarization process which examines the evidential structure supporting a hypothesis extension. *Alternative extension* SOUs result in the evaluation of the alternative hypothesis extensions using the same process. The evaluation functions effectively compute the conditional probabilities for the hypothesis extensions. Domain-specific evaluation functions are currently used because neither Bayes’ Rule nor Dempster’s Rule are applicable to interpretation in general because interpretation evidence typically fails to meet the necessary independence

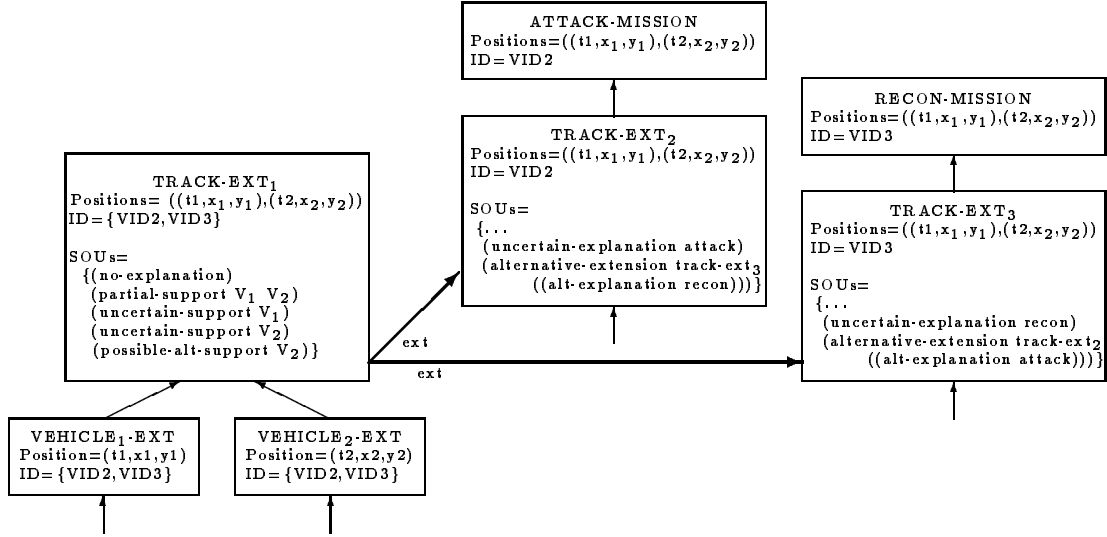


Figure 2: Example hypothesis extensions with their symbolic sources of uncertainty.

criteria. Despite this fact, the use of composite ratings does permit the use of modular evaluation functions (see [Pearl 1988]).

Interpretation vs. Classification

In order to understand why we have taken the approach that we have, it is necessary to recognize the differences between interpretation problems and similar problems that we will refer to as “classification problems.” Clancey [6] has distinguished between two basic types of problem solving approaches: *classification problem solving* and *constructive problem solving*. In classification problem solving, solutions are selected from among a *pre-enumerated* set of possible solutions. “Classification problems” are those problems that can be solved using classification problem solving techniques. This includes many of the kinds of diagnosis problems that have been studied—e.g., [1, 12]. There are some well-developed numeric techniques (e.g., Bayesian networks [11]) that are applicable to classification problems.

Constructive problem solving is required when the set of possible solutions cannot be pre-enumerated; the possible solutions for the problem must be determined as part of the problem solving process. In general, interpretation problems require constructive problem solving because the combinatorics of their answer spaces preclude the enumeration of all potential solutions. Thus while interpretation involves the classification of data, interpretation problems cannot be solved using the techniques that suffice for classification problems. Clancey [6] notes that constructive problem solving systems must be able to incrementally create and extend hypotheses, maintain large numbers of incomplete hypotheses, and apply significant amounts

of knowledge to focus the construction process. In other words, interpretation problems require search so control is critical for interpretation systems.

There are a variety of factors responsible for interpretation problem answer spaces being very large. One reason is that there may be an extremely large or even infinite number of possible hypotheses that must be considered; interpretation hypotheses are compound structures that include parameters which typically have either large numbers of discrete values or continuous values. As a result, there are often a very large or even infinite number of possible *versions* of each type of hypothesis. For example, Track hypotheses in an aircraft monitoring system include an “ID” parameter that represents the type of the aircraft out of all of the aircraft types and a “positions” parameter that represents the X-Y positions of the aircraft over time. The set of possible versions of any Track hypothesis is then the “cross-product” of the sets of possible values for each of these parameters. Even if positions are not represented by continuous values (because sensor resolution limitations are taken into account), there can be an extremely large set of possible Track hypotheses.

The combinatorics of representing hypotheses is further complicated by the fact that evidence for a hypothesis may only partially constrain the values of the hypothesis’ parameters. In other words, evidence for a hypothesis may leave the hypothesis’ parameters uncertain—it may support a subset or range of the possible values for each hypothesis parameter. For example, each piece of acoustic sensor data may be capable of supporting Track hypotheses with a number of different aircraft IDs because several different types of aircraft can produce the same acoustic frequencies. Furthermore, limitations in sensor resolution results in uncertainty in the ac-

tual acoustic frequency being sensed which leads in turn to an even greater number of aircraft types that might be supported by each piece of acoustic sensor data. It is only by combining many pieces of sensor data that this aircraft ID uncertainty can be resolved—i.e., that the subset of possible aircraft types which the available data supports can be determined.

Another source of solution combinatorics for interpretation problems is the possibility of multiple correct hypotheses of each type (i.e., multiple correct instances of each interpretation type). For example, in an aircraft monitoring application, multiple aircraft may be monitored so there may be multiple correct Track (or Mission-level) hypotheses. However, the number of aircraft that will be monitored in a given region and period of time cannot be known a priori. This means that the set of potential “solutions” to the aircraft monitoring problem must include no Track hypotheses, one of any of the possible Track hypothesis versions, two of any of the possible Track hypothesis versions, etc.—up to some maximum number of vehicles that might be monitored (in the specified region and time).

Another consequence of the possibility of multiple instances of hypotheses is that the goal of an interpretation system is somewhat different from that of a classification system. The overall goal of an interpretation system is not only to resolve its uncertainty about the correctness of the hypotheses it has created, but to be sure that these hypotheses cover all of the valid interpretations. For example, an aircraft monitoring system must not only resolve uncertainty about the correctness of any Track or mission-level hypotheses that it creates, but must also be sure that it has examined enough of the data to create hypotheses for all possible aircraft.

The possibility of multiple correct hypotheses of each type also results in what is known in data fusion terminology as the problem of “correlation ambiguity” [8]. What this means is that even if it is certain that a hypothesis supports (is explained by) some *type* of hypothesis, it may still be uncertain which *particular* hypothesis (of that type) it supports. Correlation ambiguity is an important source of uncertainty for interpretation problems, but does not affect classification problems.

Another way to look at the differences between interpretation and classification problems is by thinking about the use of belief networks [11]. Classification problems can be approached using belief networks (though there may still be some difficult problems involved in determining the “best” answer in terms of possible covering sets [12]). Interpretation problems cannot be approached in this way: it would be difficult or impossible to instantiate a network of all possible hypotheses, this would also be very inefficient since only a small percentage of the possibilities will be supported by the data, be-

cause an indeterminate number of instances of each type of hypothesis may be correct the system effectively needs to instantiate an indeterminate number of networks, and the network to which any piece of data applies would be uncertain. When interpretation problems have been “solved” with classification techniques alone (see [9, 10]), what has actually been done is that many of the difficult aspects of the problems have been simplified or ignored. See [2, 3] for further discussions of these issues.

Related Research

Numeric representations of uncertainty like probabilities and Dempster-Shafer belief functions cannot be used to identify methods for directly resolving uncertainties because they *summarize* the reasons that the evidence is uncertain [11]. We have chosen to maintain a representation of the reasons that hypotheses are uncertain in order to allow the system to use a wide range of strategies to resolve uncertainty. Our use of a symbolic representation of uncertainty is similar to Cohen’s [7] symbolic representations of the reasons to believe and disbelieve evidence which he calls *endorsements*. However, whereas Cohen was trying to develop a semantics for general evidential reasoning, our representation is tailored to the needs of interpretation control. This means that we only had to be concerned with the one type of well-understood inference which is the basis for interpretation: abductive inference. Cohen’s task was very difficult because he was trying to capture the uncertainty in a wide variety of poorly understood types of inferences and develop methods for combining the endorsements from these inferences. Because of the complexity of the problem, the work on endorsements did not result in any general-purpose formalism for representing and reasoning with symbolic uncertainties. Our work demonstrates that it is possible to maintain and reason with detailed information about the sources of uncertainty in evidence when dealing with well-defined types of inferences.

Since an interpretation specification is effectively a specification of causal relations, the network of interpretation hypotheses that is constructed by the interpretation process is similar to a Bayesian network [11]. However, because it is impossible to pre-enumerate the possible solutions, interpretation is not simply a matter of instantiating a belief network and then propagating probability information as new evidence is added. It is also important to recognize that while Pearl’s work addresses the issue of evaluating belief given a set of evidence, it does not address the problem of identifying evidence which could be gathered to resolve uncertainty—unless nodes for such evidence are already instantiated in the network (without any “diagnostic evidence”). Thus Pearl’s work does not eliminate the need for explicit representations of the factors affecting belief

if one is to decide *how* to resolve uncertainty—i.e., if one is to make control decisions (our principal focus). Our system does include one of the key features of the Bayesian network formalism: the use of the causal relationships to recognize information relevant to a hypothesis.

Conclusions

In order to evaluate this framework, we have implemented the concepts with a simulated aircraft monitoring application. Aircraft monitoring is a suitable domain for the evaluation because it has characteristics that exercise all of the capabilities of the system: there are multiple sources of evidence from multiple types of sensors, some of these are active sensor which are under the control of the system, there are complex interactions between competing hypotheses, and there are large numbers of potential interpretations of the data due to the modeling of ghosting, noise, and sensor errors. The experiments that have been run (see [3, 4]) have demonstrated that this framework can support a range of methods for resolving uncertainty and that these methods can improve the performance of blackboard-based interpretation systems.

To confirm the generality of the model of interpretation uncertainty, we have also implemented a testbed for sound understanding in household environments—such as would be necessary for robots. So far, this research has shown that the set of SOU classes is sufficient, but that specific SOUs instances may need to be adapted for particular applications in order to enhance control decisions. One issue that we are pursuing with respect to sound understanding is how low-level processing can be best handled even when it doesn't fit into an abductive framework. A multi-agent, distributed version of the aircraft monitoring application is being pursued to evaluate the usefulness of the SOU representation for driving the communication among agents and for doing distributed differential diagnosis [5].

References

- [1] Buchanan, Bruce and Edward Shortliffe, editors, *Rule-Based Expert Systems*, Addison-Wesley, 1984.
- [2] Carver, Norman, *Control for Interpretation: Planning to Resolve Uncertainty*, Technical Report 90-53, Computer and Information Science Department, University of Massachusetts, 1990.
- [3] Carver, Norman, *Sophisticated Control for Interpretation: Planning to Resolve Uncertainty*, Ph.D. Thesis, Computer and Information Science Department, University of Massachusetts, 1990.
- [4] Carver, Norman and Victor Lesser, "A New Framework for Sensor Interpretation: Planning to Resolve Sources of Uncertainty," to appear in *Proceedings of AAAI-91*, 1991.
- [5] Carver, Norman and Victor Lesser, "Sophisticated Cooperation in FA/C Distributed Problem Solving Systems," to appear in *Proceedings of AAAI-91*, 1991.
- [6] Clancey, William, "Heuristic Classification," *Artificial Intelligence*, vol. 27, 289–350, 1985.
- [7] Cohen, Paul, *Heuristic Reasoning About Uncertainty: An Artificial Intelligence Approach*, Pitman, 1985.
- [8] Lakin, W.L., J.A.H. Miles, and C.D. Byrne, "Intelligent Data Fusion for Naval Command and Control," in *Blackboard Systems*, Robert Englemore and Tony Morgan, editors, Addison-Wesley, 1988.
- [9] Lowrance, John and Thomas Garvey, *Evidential Reasoning: An Implementation for Multi-sensor Integration*, Technical Note 307, SRI International, 1983.
- [10] Lowrance, John, Thomas Garvey, and Thomas Strat, "A Framework for Evidential-Reasoning Systems," *Proceedings of AAAI-86*, 896–903, 1986.
- [11] Pearl, Judea, *Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Inference*, Morgan Kaufman, 1988.
- [12] Peng, Yun and James Reggia, "Plausibility of Diagnostic Hypotheses: The Nature of Simplicity," *Proceedings of AAAI-86*, 140–145, 1986.