

Ontology Summit 2019 Communiqué: Explanations

Kenneth Baclawski¹, Mike Bennett², Gary Berg-Cross³,
Donna Fritzsche⁴, Ravi Sharma⁵, Janet Singer⁶, John Sowa⁷,
Ram D. Sriram⁸, Mark Underwood⁹, David Whitten¹⁰

¹ Northeastern University, Boston, MA USA, ²Hypercube Limited, London, UK,
³RDA US Advisory Group, Troy, NY USA, ⁴Hummingbird Design, Chicago, IL USA,
⁵Senior Enterprise Architect, Elk Grove, CA USA, ⁶INCOSE, Scotts Valley, CA USA,
⁷Kyndi, Inc. USA, ⁸National Institute of Standards & Technology, Gaithersburg, MD USA,
⁹Krypton Brothers LLC, ¹⁰WorldVistA USA

Abstract

With the increasing amount of software devoted to industrial automation and process control, it is becoming more important than ever for systems to be able to explain their behavior. In some domains, such as financial services, explainability is mandated by law. In spite of this, explanation today is largely handled in an unsystematic manner, if it is handled at all. The decisions of modern AI systems, such as those built on deep neural networks, are especially difficult to explain. The goal of the recent Ontology Summit 2019 was concerned with the role of ontologies for explaining the functioning of a system. More specifically, the Ontology Summit focused on critical explanation gaps and the role of ontologies for dealing with these gaps. The sessions examined current technologies and real needs driven by risks and requirements to meet legal or other standards. The sessions covered explainable artificial intelligence, commonsense reasoning and knowledge, the role of narrative, and explanations in the fields of finance and medicine. The goal of this Communiqué is to foster research and development of approaches to explanations and to drive towards explanation support which can be incorporated into both knowledge engineering processes and ontology design best practices.

Keywords: ontologies; context; inference; commonsense reasoning; narrative; explainable AI; financial explanations; medical explanations

1 Introduction

The last few years have seen a substantial increase in the use of machine learning (ML) techniques for solving important problems in the field of Artificial Intelligence (AI). These advances have been driven by the availability of massive amounts of raw data. The ML techniques construct complex statistical models by processing large datasets. Unfortunately, it is difficult to explain how these ML models come to their conclusions, since each decision is, in principle, the result of a program that includes the entire

dataset that was used to develop the model. This has led to the perception that ML models are too “deep” and “mysterious” to be adequately explained. Whether or not this is an accurate perception, to solve the problem of explaining the functioning of an ML model or, indeed, any large complex system, it is necessary to address the issue of what an explanation is, as well as what criteria can be used to evaluate the accuracy of an explanation and its suitability for a purpose.

In his presentation at the Ontology Summit 2019,[9] Derek Doran asked the question “Okay but Really... What is Explainable AI?” He found that there was wide disagreement about the answer to this question. As a first approximation to classifying this variety, he identified a hierarchy of explanation “types” an AI system can exhibit:

1. Interpretable: I can identify why an input goes to an output.
2. Explainable: I can explain how inputs are mapped to outputs.
3. Reasoning: I can conceptualize how inputs are mapped to outputs.

To illustrate the distinctions among the three levels, consider the case of an application for a loan that was denied. The following are examples of the three types of explanation for answering the question “Why was my application for a loan denied?”

1. Interpretation: The account balance covariate in the logistic regression model used to make decisions explains 89.9% of residual variance.
2. Explanation: The system denies loan applicants with low bank account balances.
3. Reasoning: The system does not want to approve loans to those who do not show evidence of being able to pay them off.

While all three types of explanation are expressed as natural language narratives, there is a significant distinction among them. The first type rigorously expresses features of the decision in terms of a statistical model. It has the most information and even qualifies as a “proof,” but fails to relate the decision to the context in which the decision was made. The second type is much better because it is now expressed in terms of the customer context, even if it is not a fully rigorous proof. Yet the second type only explains the function without explaining why that function is being used. The third type is the best since it now explains why that function was used by the bank. The third type uses formal reasoning to infer the rationale from the other types of explanation (i.e., Interpretation and Explanation). What is not shown in the example for the Reasoning type is that it should allow for an interaction with the customer with the goal of achieving customer acceptance of the explanation.

While the first two types above may be adequate in some limited technical contexts, for the most part, humans would demand the third and highest level. Accordingly, there was general agreement with the following definition of explanation, “An explanation is the ability to answer a how or why question and to answer a follow-up question to clarify in a particular context.” This definition makes it clear that an explanation can never be enough if it consists of a single generated answer, no matter how elaborate it might be. There must be a capability for a dialog between the user and the system. A further consequence of this definition is the importance of context. Since the notion of context was the subject of the Ontology Summit 2018,[4] the current Summit may be regarded as a continuation of the previous Summit. Indeed, since context includes the answers to all of the “six Ws” (namely, Who, What, When, Where, Why and

How¹), explanation is both part of the context and depends on it.

The Ontology Summit 2019 sought to explore, identify and articulate how ontology can bring value to the problem of automating explanations of complex systems in general. The Summit dealt with this goal by first studying the notion of explanation in a series of sessions in the Fall of 2018. The two most important areas that were identified at this time were (1) Commonsense Reasoning and Knowledge (CSRK), and (2) Narratives. These were then studied at greater depth in subsequent sessions in 2019. In addition, it was decided that some specific domains should be studied to determine the kinds of problems that practitioners are facing with respect to explanations. The two chosen domains were the Financial and the Medical domains, both of which were broadly defined. Finally, since explanations for AI were the original motivation, the problem of Explainable AI was also studied. Accordingly, the Summit consisted of sessions devoted to five tracks: Explainable AI, Commonsense Reasoning and Knowledge, The Role of Narrative, Financial Explanations, and Medical Explanations.

The purpose of this Communiqué is to identify some of the prevailing viewpoints and the major issues and challenges of explanation, especially the role that ontology could play toward solving these challenges. The Communiqué begins with some background and history of the notion of explanation in the next section. This is followed by sections devoted to synthesizing the findings from each of the five tracks in Sections 3 to 7. The Communiqué then summarizes the findings in Section 8. We end with a conclusion in Section 9 and acknowledgments in Section 10.

2 Background

An explanation is the answer to the question “Why?” as well as the answers to follow-up questions such as “Where do I go from here?” Accordingly, explanations generally occur within the context of a process, which could be a dialog between a person and a system or could be an agent-to-agent communication process between two systems. Explanations also occur in social interactions when clarifying a point, expounding a view, or interpreting behavior. In all such circumstances in common parlance one is giving/offering an explanation.

Among the first known attempts at understanding the ‘why’ of explanations were those documented among Indian intellectuals and philosophers, beginning with the knowledge collection called the Vedas (dating back to 5000 BCE). This philosophical tradition included notions of context, logic and explanation that are similar to the modern conceptions. For example, there was a notion of syllogism that explicitly incorporated context into the structure of the syllogism. Explanation was also a part of logical inference. More generally, explanation in the form of a dialog between a teacher and a student appears throughout the Vedas.

Later in the classical Greek period, Aristotle provided a view of explanation as part of a deductive process using reason to reach conclusions. In the Posterior Analytics from his *Organon*, Aristotle proposed 4 types of causes (*αἰτίαι*) to explain things. These were from either the thing’s matter (material cause), form (formal cause), end (purpose or final cause), or agent (efficient cause).

The Ontology Summit 2019 was concerned with the role of ontologies for explaining the reasoning of a system. More specifically, the Summit focused on critical explanation gaps and the role of ontologies for

¹These are called the “six Ws” in spite of one of them not starting with W.

dealing with these gaps. The summit sessions used current technology to review these gaps and roles. Real needs driven by risks highlighted issues beyond the theoretical ones .

Inspired by the current DARPA Explainable AI (XAI) Program,[15, 16] the Ontology Summit theme considered the general problem of explanation. The Summit considered not only AI systems that can explain their actions, but also other smart engineering systems which cooperate with each other and aid humans. With the increasing amount of software devoted to industrial automation and process control, this capability is becoming more important than ever. Explanations include expressing rationales, characterizing strengths and weaknesses, and projecting their behavior into the future.

Ontologies could play a significant role in explanation since smart engineering systems must represent the conceptual framework that supports explanation. An ontology used for explanation would include terms for domain and natural world concepts, relations, and activities. Some version of natural language may be used to describe states and actions in terms that people easily understand. This system could also provide the conceptual structures for the system’s dialogs, plans and actions used in explanation.

A benefit of the use of ontologies in support of explanations is the potential for improving interoperability between systems that otherwise would not have a common framework for interoperation. The danger is that current efforts for explainability will have many of the same undesirable characteristics that have been criticized AI systems. One of these is fragility (also called brittleness) in the following sense: two perceptively indistinguishable inputs with the same predicted label can be assigned very different interpretations and explanations.[11] Another is siloing in the sense that each product employs completely different and incompatible explanation techniques. While such products individually satisfy the requirement of providing explanations, they cannot be consistently integrated into large scale systems.

3 Explainable Artificial Intelligence

Since the 1950s, when the term Artificial Intelligence was coined, there has been considerable progress in this area. The 1980s was dominated with the rise of rule-based systems. This period has been called “the first wave” of AI. Advancements in computer hardware facilitated multilayered neural networks, which led to significant improvements in machine learning for certain classes of problems in the 2000s. This was the “second wave.” Now, we are witnessing the “third wave,” which will include a combination of neural networks and knowledge structures; DARPA views the third wave as contextual adaptation where “systems construct contextual explanatory models for classes of real world problems.”

The first wave of AI produced several reasoning techniques, such as decision trees, inference networks, Bayesian networks, statistical learning techniques, and the beginnings of neural networks. During the first wave, it was already recognized that these systems were unable to provide adequate explanations.[18] As Ford et al surmised, “...the inadequate approach to explanation found in most expert systems can be a significant impediment to user acceptance.”[10]

As the second wave took place, multilayered neural networks (deep learning) with more than one hidden layer produced impressive results in a wide variety of classification problems. There was a tradeoff between performance and explainability. For example, a major problem with neural networks is the lack of transparency. It was more like a black box approach, where the following questions could not be adequately answered:

- Why did you do that?
- Why not something else?
- When do you succeed?
- When do you fail?
- When can I trust you?
- How do I correct an error?

The DARPA XAI Program initially had a goal of addressing the above questions by enhancing AI software to deal with such questions.[16] See Figures 1 and 2. Note that the DARPA XAI Program proposed that ontologies would be part of the solution. The program also specified that counterfactual arguments are required.

Explainable AI defines a field of research, not a particular “type” of AI. There exist unique notions of ‘explanations’, making problem- or context-independent explainable AI work inappropriate. For AI systems to be understandable, symbolic systems and reasoning must be integrated to provide operationally effective communications of the internal state and operation of the system.[22] Meaningful ontologies and knowledge bases, and Symbolic AI methods are fundamental (but currently, mostly ignored) to the most important explainable AI use-case: when in practice, within some context, a layman must understand, trust, and be responsible for the conclusions an AI system draws.

Sargur Srihari divided the history of AI into three waves: (1) classic rule-based machine learning, (2) deep learning, and the next wave, (3) probabilistic AI. While probability and statistics are related, they are fundamental differences between them. Statistics analyzes observations to understand them, most commonly in the form of a probabilistic model. Probability deals with models that allow one to make predictions and decisions when there is uncertainty. The probabilistic models can be postulated *a priori* or estimated *a posteriori* using statistics. The next wave of AI that Sargur Srihari proposed will focus on generative models in the form of probability distributions that can be queried. The algorithms of this wave will use a variety of algorithms and architectures for knowledge representation, inference and learning. It will use Probabilistic Graphical Models (PGMs) to find the most probable explanation (MPE), also called the maximum a posteriori probability assignment. As a result, the third wave may be more suited to explainability than the first two waves. He has used PGMs and MPE in his work on forensics which lends itself to a combination of deep learning and probabilistic explanation.[28]

4 Commonsense Reasoning and Knowledge

There is a long history showing the relevance of commonsense reasoning and knowledge (CSRK) to explanation. Certainly many pioneers of modern AI believed so and argued that a major long-term goal of AI should include endowing computers with standard commonsense reasoning capabilities. In “Programs with Common Sense”[21], John McCarthy described 3 ways for AI to proceed:

1. imitate the human central nervous system,
2. study human cognition or
3. “understand the common sense world in which people achieve their goals.”

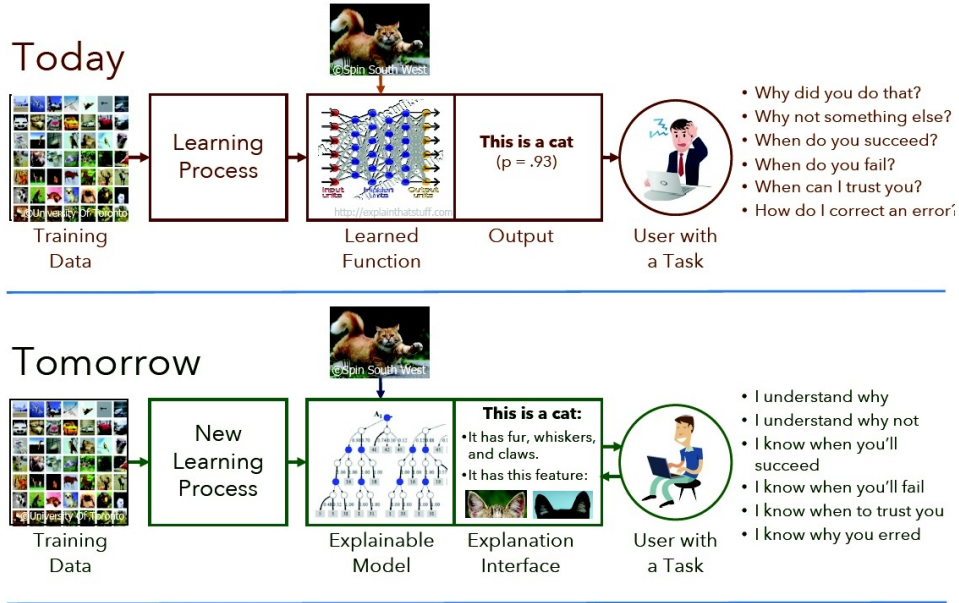


Figure 1: The DARPA XAI Program [16]

Another related goal has been to endow AI systems with natural language (NL) understanding and production. We see the relation of CSRK and explanation in the early conceptualization of a smart advice taker system from McCarthy's work that would have causal knowledge available to it for: "a fairly wide class of immediate logical consequences of anything it is told and its previous knowledge." McCarthy further noted that this useful property, if designed well, would be expected to have much in common with what makes us describe certain humans as "having common sense." He went on to use this idea to define CSRK – "We shall therefore say that a program has common sense if it automatically deduces for itself a sufficiently wide class of immediate consequences of anything it is told and what it already knows." [21]

There is a long path of research from these early days to today which includes simple logical traces of rule firings which proved to be brittle and did not capture deep knowledge. Early AI research demonstrated that the nature and scale of the both explanation and CSRK was difficult. People seemed to need a vast store of everyday knowledge for common tasks which they also used as common understanding and explanation. A variety of knowledge was needed to understand even the simplest children's story which is a feat that children master with what seems a natural process. One resulting approach was an effort like CyC to encode a broad range of human commonsense knowledge as a step to understanding text which would bootstrap further learning. But this type of progress has been slow and some believe that today this problem of scale can be addressed in new ways including via modern machine learning. But these methods do not build-in an obvious way to provide machine generated explanations of what they "know." For this reason research is exploring how common sense could assist in addressing some challenges in Deep Learning (DL). Particular challenges include addressing the known brittleness of DL models given various adversarial changes to input, and generalization to unseen situations when faced with limited training data. Preliminary work suggests that CSRK can compensate for limited training data and make it easier to generate explanations, assuming that common sense is available in an easily consumable representation.

Performance vs. Explainability: DARPA XAI Program

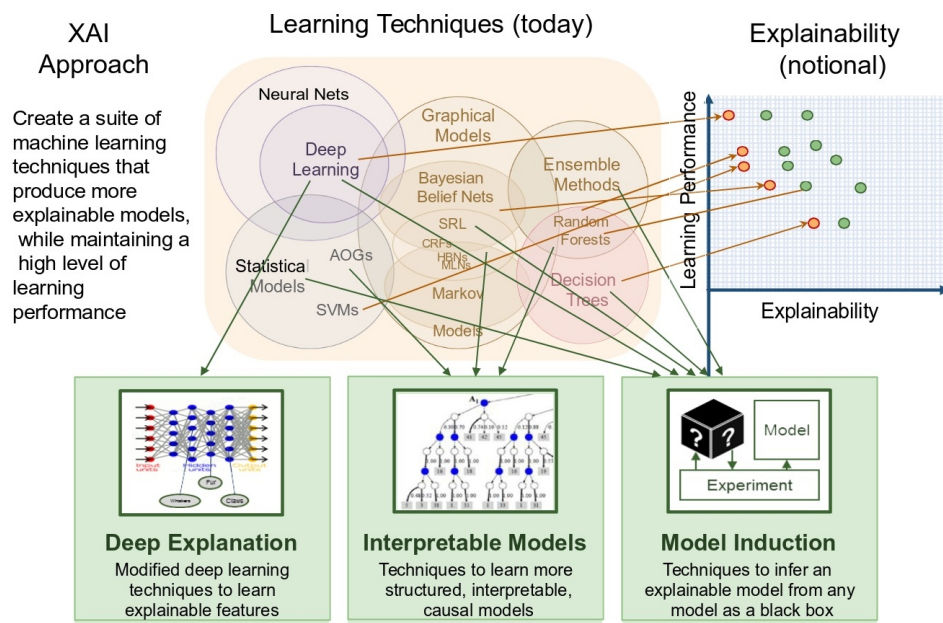


Figure 2: Performance vs Explainability [16]

There is evidence, for example, that state-change commonsense knowledge can be used for limited training data for procedural text to help prediction. Common sense aware DL models makes more sensible predictions despite limited training data using the prior knowledge of state-changes, e.g., it is unlikely that a ball gets destroyed in a basketball game scenario. The model injects common sense at decoding phase by re-scoring the search space such that the probability mass is driven away from unlikely situations. Adding common sense results in much better and more robust performance.

Explanations in the DL context can be discussed along various dimensions such whether they can provide rationales to justify decisions, or whether they can integrate multimodal information. In recent years, datasets have fueled research in DL. In both the NLP and computer vision communities, datasets containing explanations have gained interest. This has led to models that provide explanation in the form of attention, natural language sentences and structures such as scene graphs and state-change matrices. Evaluating explanation has remained a challenge, and some evaluation metrics include string matching, human based judgments, automated evaluations, and, finally, learning to evaluate. CSRK is helpful along these metric areas and while an important asset for DL models overall, its logical representation in such forms as microtheories has not been successfully employed. Instead, tuples or graphs consisting of natural language nodes has shown some promise, but these face the problem of string matching when there are linguistic variations. More recent work on supplying common sense in the form of adversarial examples, or in the form of unstructured paragraphs or sentences has been gaining attention. Also of note is the fact that CSRK could prove valuable as an aid to extended AI system-human dialog since a good deal of CSRK is involved in keeping track of conversational threads.

Some general recurring questions that are worth considering include:

1. The two most common approaches to CSRK are (1) formal representations (e.g., encoded in an ontology’s content as in Cyc or Pat Hayes’ Ontology of Liquids)[17] and (2) strategies for common-sense reasoning (e.g., default reasoning, prototypes, uncertainty quantification, etc.) How can one leverage the best of these two approaches?
2. How can one inject commonsense knowledge into machine learning approaches? There has been some progress on learning using taxonomic labels, but it just scratches the surface
3. How can one bridge formal knowledge representations (formal concepts and relations as axiomatized in logic) and representations of language use (e.g., Wordnet)?

The conclusion is that CSRK could assist in addressing some of the challenges of explainability in general and for explainability of DL systems in particular.

5 The Role of Narrative

Narrative may be regarded as being complementary to commonsense reasoning and knowledge. While CSRK provides the underlying reasoning and knowledge necessary for explanation, the topic of narrative covers 1) the use of knowledge elements in forming stories as sequences of events, and 2) the presentation of these stories in dialog.

Narrative is a strategy for sense-making. The Project Narrative at Ohio State University states, “Narrative theory starts from the assumption that narrative is a basic human strategy for coming to terms with fundamental elements of our experience, such as time, process, and change, and it proceeds from this assumption to study the distinctive nature of narrative and its various structures, elements, uses, and effects.”[24]

As already noted above in the very definition of explanation, an explanation narrative is an interactive conversation. People leverage their partial understandings of the causal structure of the world “by knowing how to access additional explanatory knowledge in other minds and by being particularly adept at using situational support to build explanations on the fly in real time.”[19]

Results from social science regarding how humans explain to each other can serve as a useful starting point for explanation in artificial intelligence.[23, 24] Key insights include:

- Explanations of social behavior rely on theories of an actor’s beliefs, desires, intentions and traits, all of which must be considered when generating explanation narratives.
- In addition to causal explanations, i.e., “Why P,” explanation narrative also requires contrastive (counterfactual) answers, i.e., “Why P rather than Q?”
- The selection of such causes tends not to be comprehensive but rather identifies a few causes considered to be adequate for the explanation.
- Similarly, one must be selective in the choices of the counterfactual simulations. The use of counterfactual explanation in narrative is a common part of conversation. Automated explanation must be capable of arguing both in favor of an answer as well as against proposed alternative answers.
- As interactive conversations, explanation narratives must abide by rules and maxims of discourse, such as those identified by Grice[12] as follows:

- The Maxim of Quality
- The Maxim of Quantity
- The Maxim of Relation
- The Maxim of Manner

Relating the conversational evaluation criteria above to the relevance, usefulness and trustworthiness criteria previously noted for explanations can yield additional insights. For example, the Maxim of Quantity suggests that an explanation should be as extensive as necessary but no more so. Surprisingly, the trustworthiness criterion is another one that also has this property. If humans trust a system too much, then they will be unprepared for situations in which the system exhibits anomalous or dangerous behavior.[25]

Many techniques have been developed for generating NL from knowledge and reasoning processes. It is commonly assumed that proofs are the “gold standard” for explanation. However, in practice, proofs are not adequate as explanations, as we remarked in Section 1 above. Surprisingly, even proofs in the Mathematics literature are far from being complete and rigorous. Their purpose is to convince other mathematicians in the same specialty (which could be quite small) that the theorem is probably true. In most cases, one could generate a complete and rigorous proof from the published argument, but it is very difficult to accomplish this and it is rarely done. In some cases, the claimed theorem turns out to be incorrect in spite of having a thoroughly reviewed and published “proof.” Alternatively, Baclawski has proposed that by taking explanation as the goal, one can develop fully complete and rigorous proofs that could also be readable and convincing narratives.[3]

Tiddi[32] studied theories and models of explanation in a variety of disciplines to identify the minimal ‘story’ elements that define an explanation. The resulting Explanation Ontology Design Pattern is shown in Figure 3. In this pattern, an explanation requires:

- An explanans (the explanation)
- An explanandum (what is being explained)
- A context (or situation) that relates the two
- Some theory or “a description that represents a set of assumptions for describing something, usually general. Scientific, philosophical, and commonsense theories can be included here.”
- An agent producing the explanation

The theory element is to be taken as including scientific, philosophical, and commonsense theories as well as concepts defined in cognitive science such as laws, human capacities, human experiences, universals. Theory is a specialization of the Description class from the Descriptions & Situations ontology. “In this view, our theory acts as the description that classifies the explanation (corresponding to the situation), which is built upon the context, the event to be explained, and the possible explaining events.”[32]

6 Financial Explanations

The financial services industry is diverse, international and huge. Depending on how it is defined, it globally represents a \$10 to \$20 trillion industry (in US currency). It is also rapidly evolving, with new

Explanation Ontology Design Pattern

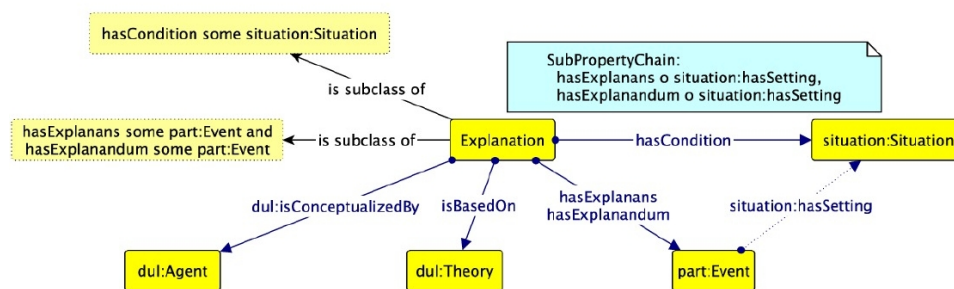


Figure 2: The Explanation Ontology Design Pattern.

Tiddi, et al. (2015)

Figure 3: Explanation Design Pattern from [32]

investment vehicles and services constantly being introduced. Although this industry is highly regulated, there is no common industry language for this domain. The US Treasury Office of Financial Research stated, “The lack of a common industry language is a billion dollar problem.” This situation is remarkable considering that business and finance were some of the earliest domains to make heavy use of computer technology.

Formalizing the financial services industry is a formidable problem. Financial regulations are very large and written in unstructured, natural language text. A single US regulation, the Federal Reserve System Final Rule 12 CFR Part 223, consists of 143 pages of text, and the summary alone is 19 pages long. Another complication is that financial regulations have long reference chains to other regulations. Nevertheless, progress is being made with the development of the Financial Industry Business Ontology (FIBO).

Financial service explanations are increasingly being mandated. Explanations must be provided to customers (which includes both individuals and companies), to financial services partners, and to government regulators. Ontologies such as FIBO can contribute to explanations but at the cost of adding the need to explain the ontology itself.

Consider the example of the credit industry, a significant subdomain of financial services. In the US, the Fair Credit Reporting Act (FCRA) mandates that credit decisions must be explained. For example, a customer whose application for a credit card might ask “Why am I being denied a credit card?” The credit card company might then respond “Because you didn’t meet our criteria.” Needless to say, the customer is likely to ask which criteria were not met, and a dialog will then be necessary.

There are many challenges involved in providing explanations in the credit industry in particular, and the financial services industry in general.

- Context
 - “Who” and “Where”: Must comply with Know Your Customer (KYC) regulations. These regulations are intended to prevent illegal activities such as money laundering and bribery.
 - “Who”: Privacy must be respected.
 - “Who”: Certain individuals are protected, such as veterans.

- “When”: There are temporal constraints, which are complicated by reporting lag times.
- “What”: Product-specific credit scenarios affect the explanations.
- Data Fusion
 - Must be coordinated with credit reporting agencies (third party).
 - Credit cards can have multiple authorized persons.
- Decision support algorithms
 - Rule-based decision support software should be explainable but could have been supplied by a different company.
 - Machine learning algorithm results are difficult to explain, even if the algorithm was applied by the same company.
- Explanations of ontology
 - Reasoning with ontologies adds a further aspect that itself needs to be explained.
 - Higher-order concepts must be explained, such as rules and mathematical constructs.
- Narrative and dialog
 - Most human understanding is contextual and this needs to be taken into account in presenting explanatory material.
 - The presentation should be easily understandable to humans.
 - Follow up inquiries for more detailed explanations should be supported.

7 Medical Explanations

The health care enterprise involves many different stakeholders – consumers, health care professionals and providers, researchers, and insurers. AI will play an important role in many tasks that these stakeholders undertake. These include: image diagnostics, medical decision making, prior authorization, drug design, nutrition advisor, patient scheduling, etc. During the late 1970s and early 1980s, the pioneering knowledge-based expert systems (KBES) were in the domain of medical diagnosis (e.g., MYCIN, INTERNIST). These systems relied on rule-based inferencing and probabilistic knowledge networks. Unfortunately, as noted in Section 3, most of these systems did not provide adequate explanations of their decisions. In 1993, Ford et al surmised, “...the inadequate approach to explanation found in most expert systems can be a significant impediment to user acceptance.”[10]

Currently, neural networks (or deep neural networks [DNNs]) are playing a key role in many medical applications that involves image interpretation and diagnosis. Eric Topol’s paper and book provide numerous examples of such DNN-based systems.[33, 34] However, the key problem with DNN-based systems is the lack of explainability, which is very important for many reasons. Medical AI systems will be accepted only if they are trusted by clinicians and other stakeholders, and explainability is essential for this, since patient lives depend on medical decisions. There are also laws and regulations giving patients the right to explanations about their diagnoses and treatments. Explanations are more significant in some medical fields, such as mental health.[20]

A combination of DNNs and knowledge graphs (ontologies) could potentially form the basis for future AI systems in medicine. Achieving this requires solving several problems as discussed in the Introduction

above. One must first be able to interpret the DNN. In other words, identify why an input goes to an output. Next one must integrate the interpretation of the DNN with the knowledge graph so that one can trace the reason for the input producing the output.[20] One must then understand the concepts in the knowledge graph. Finally, a human readable narrative must be generated as discussed in Section 5.

While there is general agreement on the steps required for the generation of explanations given above, there are several possibilities for how knowledge graphs could be used in the machine learning architecture. While there are many ML algorithms, most are variations on supervised learning, i.e., first train the model on labeled training data, then apply the model to new unlabeled data. Based on this architecture, there are at three ways in which knowledge graphs could be used for explainability:

1. Knowledge enhancement *before* a model is trained
2. Knowledge harnessing *after* the model is trained
3. Knowledge infusing *while* the model is employed

Kursuncu et al have developed and evaluated all three of these roles for knowledge graphs in the mental health domain using the same ontology for each role.[20]

Other forms of learning, such as semi-supervised and unsupervised learning may also have multiple roles for ontologies. The Global Health Intelligence project, which involved data integration and interoperability for malaria monitoring and evaluation in sub-Saharan Africa, integrated and analyzed large volumes of structured and unstructured data, and processed the data to find patterns. The project also produced explanatory and predictive models.[26] In this case, each of the steps, integration, analysis, pattern recognition and model development can make use of ontologies prior to, during and following the step.

8 Findings

We now summarize the main findings of the Summit. The findings were organized into groups that very roughly correspond to the phases of a system development project. Ontologies can play important roles in achieving explainability throughout all development phases. We begin in Section 8.1 with the reasons *why* a system should be able to explain its functioning. This is followed by Section 8.2 which proposes *what* is necessary for explainability. The last two sections discuss *how* explainability could be achieved. Section 8.3 surveys some potential approaches, and Section 8.4 lists some of the challenges that would need to be overcome. Each challenge represents an opportunity for research on explainable systems.

8.1 Rationale

As we have seen, systems that can explain themselves are difficult to develop. Accordingly, there must be a compelling rationale for allocating resources to support explainability. In this section we mention a few of the reasons why a system should be explainable.

Perhaps the most all-encompassing reason for developing explainability is the General Data Protection Regulation (GDPR) which states, “The data subject should have the right ... to obtain an explanation of the decision reached ...”[1, 31] However, the extent to which the GDPR regulations themselves provide a

“right to explanation” is heavily debated. Given that explainability is still a nascent research activity, it should not be surprising that the GDPR stops short of mandating this ability immediately. Nevertheless, the intent is to make it easier for users to obtain explanations, especially for fully automated decision making systems.

While explainability is not yet mandated in general, some domains already have legal requirements. Financial services have legal requirements for explainability.[5, 36] Forensics is another domain that effectively requires explainability since forensic reports are intended to be used in legal proceedings.[28] Medical systems often have legal requirements for explainability, as well as guidelines and protocols that must be followed.[20]

Even in domains that do not have requirements for explainability, the need is still compelling. In many domains, AI is still a novelty and explanation may be necessary to gain acceptance and trust by practitioners.[20, 35]

It is known that AI systems are fragile and can be easily fooled.[11] Explanation can help to improve robustness and to deal with adversaries.[8, 30]

8.2 Requirements

Having made a case for systems to be explainable, the question then becomes what is required. As discussed in Section 2 above, an explainable system must be able to answer commonsense questions, especially “why” questions but also “how” and “why not” questions. The questions and answers must be in natural language, must be expressed in an easily understandable narrative style, and must allow for interactivity and following up.[6, 8, 27, 29]

There was a general consensus among the invited speakers and participants that ontology is important in explainability. While ontologies have been proposed for explanation, for example, the Explanation Ontology Design Pattern[32], there does not appear to be a single predominant role that ontologies play in explanation. An ontology could be in the “background” for purposes such as commonsense reasoning.[6] On the other hand, an ontology could also be more explicitly visible to end users as in financial services.[5, 36] When this is the case, the ontology must itself be explainable.[5] When an ML algorithm is being employed, there is a great variety of roles that an ontology could play as discussed in Section 7.[20, 26] Some of these are background roles, while others are explicit and require the ontology to be explained.

It was also generally agreed that context is important.[5, 6, 8, 29, 36] Indeed, as noted in Section 1, context is part of the definition of explanation. Moreover, explanation may be regarded as being a part of the context as well as depending on it. Specifically, explainability requires a “user model” of interests and knowledge in order to have meaningful dialogs.[8] A more subtle aspect of the user model is the user’s language. This is not just the natural language of the user but also the domain-specific jargon of the user’s community.

Finally, there is a need for evaluating the quality of explanations.[2, 14] Quality includes not only subjective measures of explainability but also whether the explanations are verifiable and conform to standards.

8.3 Development

We now survey some of the most promising approaches to achieving explainability that were presented at the Summit.

A common theme during the Summit sessions was that explainability is very likely to require cross-disciplinary teams from computer science, cognitive science, linguistics and social science. Within computer science, many fields are involved, such as classical logical (symbolic) AI, statistical (sub-symbolic) AI, NL processing, NL generation, ontologies, and CSRK. The following are some of the challenges for developing explainable systems that were proposed in Summit sessions:

- There is a need to combine the classical logical with statistical approaches to AI and to make the provenance traceable.[29]
- While theory and proofs are an important part of any solution to explainability, automated proofs are not easy to understand by a layman[2, 6]. There is a need to apply linguistic and narrative theories to mathematical proofs to produce more easily understandable explanations.
- Proofs involving quantitative explanations are also generally incomprehensible to a layman.[9] New linguistic and narrative approaches are needed. However, translations of explanations between narratives and logic should be seamless and unambiguous so as to avoid any confusion for machine interpretability.
- What is the best way to inject CSRK into ML approaches?[7]
- Combining CSRK and NL understanding is a key enabler for a robust AI explanation system.[6, 30, 31] This combination is needed to deal with incomplete and uncertain information in dynamic environments responsively and appropriately.[6] Explainability is especially critical for everyday tasks.[14]
- Sharing and reusing insights from the social and cognitive sciences is a key enabler for developing explainable systems.[31]
- The Most Probable Explanation is a very promising approach for explainability.[28]
- Visual techniques are a promising approach to explanation. They can are gaining maturity and could be used both for input from and output to the user. Combining NL and visual techniques could be an especially effective approach to explainability.

8.4 Challenges and Opportunities

Explainability is a difficult research problem, but there was general agreement that it will be possible to develop systems that can produce explanations that are as informative and useful as an explanation from a well-informed person. That said, there are some serious challenges that represent opportunities for research.

By the definition of explanation, it depends on context, both the system context and the user context. Consequently, explainability will require solutions to many of the problems that were found in the Ontology Summit 2018.[4]

In the financial domain, regulations are very large, unstructured natural language documents with long reference chains. Formalizing such regulations correctly and efficiently is a barrier to automating expla-

nations in this domain. A further complication is that the constructs need not be first-order, making it difficult to apply existing reasoning tools.[5]

One of the requirements of explainability (and of strong AI) is commonsense reasoning. Unfortunately, while great progress has been made on this problem, it remains a research problem. Consequently, common sense is “perhaps the most significant barrier” between the focus of AI applications today, and the human-like systems we dream of.[6]

One barrier for developing explainable systems that interact with humans on everyday tasks is the lack of a common ontology for such tasks. Does there exist an ontology that we all share for representing and reasoning about the physical world?[14]

Software today, whether a legacy system or a newly developed system, most commonly does not include explainability. Unfortunately, explainability cannot simply be added as another module. Rather it should drive the software engineering process from the earliest stages of planning, analysis and design. Explainability needs must be empirically discovered during these stages.[8] Unfortunately, there is little sensitivity to these needs as well as little experience with addressing them in current systems and work practices.[8] When faced with legal requirements for explainability, there is a disconnect among the views of the stakeholders. Users generally view explainability as very useful, but investors do not perceive that it adds value.[13] Still other classical software engineering problems arise such as “mission creep” in which requirements are added later in the development process at a time when it is too late to modify the design to accommodate the new requirements, resulting in cost overruns, delays and poor designs.[13]

Finally, for the quality of explainability to be based on more objective measures than user satisfaction surveys, we need to develop and agree on appropriate standards.

9 Conclusion

Automating explanations of complex systems today is largely handled in an unsystematic manner, if it is handled at all. Modern AI systems are especially difficult to explain, but they are not the only systems that lack adequate explainability. We found that ontologies and context sensitivity are important parts of any effective solution to explainability. This conclusion was apparent in all of the Summit tracks, both in the general areas of commonsense reasoning and narrative and in the more specific domains of explainable AI, financial services and medicine. While we listed many potential approaches that are actively being pursued, much more research and development is needed for systems to be adequately explainable.

10 Acknowledgments

Certain commercial software systems are identified in this paper. Such identification does not imply recommendation or endorsement by the National Institute of Standards and Technology (NIST) or by the organizations of the authors or the endorsers of this Communiqué; nor does it imply that the products identified are necessarily the best available for the purpose. Further, any opinions, findings, conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of NIST or any other supporting U.S. government or corporate organizations.

We wish to acknowledge the support of the ontology community, especially the invited speakers and participants who contributed to the Ontology Summit. There were 23 invited speakers, some of whom gave presentations at more than one session. The complete list of sessions, speakers, and links to presentation slides and video recordings is available at <http://bit.ly/2WTduUB>. We also wish to acknowledge feedback from many other individuals, including Dieter Gawlick, Paul Sonderegger, Anna Chystiakova, Kenny Gross, Zhen Hua Liu, and Azamat Abdoullaev.

References

- [1] Recital 71 GDPR. Technical report, European Union, 2018. Retrieved June 14, 2019 from <http://bit.ly/2KL6LVz>.
- [2] K. Baclawski. Introduction to the Summit Tracks, January 2019. Retrieved on June 12, 2019 from <http://bit.ly/2VRt6DV>.
- [3] K. Baclawski. Proof as explanation and narrative, January 2019. Retrieved on June 12, 2019 from <http://bit.ly/2RqQJQ5>.
- [4] K. Baclawski, M. Bennett, G. Berg-Cross, C. Casanave, D. Fritzsche, J. Ring, T. Schneider, R. Sharma, J. Singer, J. Sowa, R.D. Sriram, A. Westerinen, and D. Whitten. Ontology Summit 2018 Communiqué: Contexts in Context. *Journal of Applied Ontology*, July 2018. DOI: 10.3233/AO-180200.
- [5] M. Bennett. Ontology Summit 2019 Financial Explanations Track Session 1: Financial Industry Explanations, February 2019. Retrieved on April 28, 2019 from <http://bit.ly/2RIViFQ>.
- [6] G. Berg-Cross and T. Hahmann. Overview of Commonsense Knowledge and Explanation, December 2018. Retrieved on April 28, 2019 from <http://bit.ly/2Uhbzxwh>.
- [7] G. Berg-Cross and T. Hahmann. Ontology Summit 2019 Launch Session: Commonsense Track, January 2019. Retrieved on April 28, 2019 from <http://bit.ly/2VZa5zg>.
- [8] W. Clancey. Explainable AI Past, Present, and Future: A Scientific Modeling Approach, February 2019. Retrieved on April 28, 2019 from <http://bit.ly/2Scjvo6>.
- [9] D. Doran. Okay but Really... What is Explainable AI? Notions and Conceptualizations of the Field, November 2018. Retrieved on April 28, 2019 from <http://bit.ly/2TM2141>.
- [10] K. Ford, A. Canas, and J. Coffey. Participatory Explanation. In *FLAIRS 93: Sixth Florida Artificial Intelligence Research Symposium*, pages 111–115, Ft. Lauderdale, FL, April 18-21 1993. Retrieved on June 5, 2019 from <http://bit.ly/318aUbX>.
- [11] A. Ghorbani, A. Abid, and J. Zou. Interpretation of neural networks is fragile. *Proceedings of the Thirty-Third AAAI Conference on Artificial Intelligence*, 33, 2019. arXiv:1710.10547.
- [12] H.P. Grice. Logic and conversation. In P. Cole and J.L. Morgan, editors, *Speech Acts*, pages 41–58. Academic Press, New York, 1975.
- [13] B. Grosz. An overview of explanation: Concepts, uses, and issues, January 2019. Retrieved on April 28, 2019 from <http://bit.ly/2R9iD30>.

- [14] M. Grüninger. Ontologies for the physical turing test, January 2019. Retrieved on April 28, 2019 from <http://bit.ly/2RiQSFz>.
- [15] D. Gunning. DARPA Explainable Artificial Intelligence: Program Update, 2017. Retrieved on June 4, 2019 from <http://bit.ly/2ITKwfd>.
- [16] D. Gunning. DARPA Explainable Artificial Intelligence, 2018. Retrieved on December 3, 2018 from <http://bit.ly/2s9d4pH>.
- [17] P. Hayes. Naive physics I: ontology for liquids. In *Readings in qualitative reasoning about physical systems*, pages 484–502. Morgan Kaufmann Publishers Inc., San Francisco, CA, USA, 1990.
- [18] P.J. Jackson. *Introduction to Expert Systems*. Addison-Wesley, Reading, MA, 1986.
- [19] F.C. Keil. Explanation and understanding. *Annu. Rev. Psychol.*, 57:227–254, 2006.
- [20] U. Kursuncu, M. Gaur, K. Thirunarayan, and A. Sheth. Explainability of Medical AI through Domain Knowledge, March 2019. Retrieved on June 12, 2019 from <http://bit.ly/2YuPUe3>.
- [21] J. McCarthy. Programs with common sense. In M. Minsky, editor, *Semantic Information Processing*, pages 403–418. MIT Press, 1968. Originally published in 1959.
- [22] D. Michie. Machine learning in the next five years. In *Proceedings of the Third European Working Session on Learning*, pages 107–122. Pitman, 1988.
- [23] T. Miller. Explanation in Artificial Intelligence: Insights from the Social Sciences. *Artificial Intelligence*, 2018.
- [24] T. Miller, P. Howe, and L. Sonenberg. Explainable AI: Beware of inmates running the asylum or: How I learnt to stop worrying and love the social and behavioural sciences. *arXiv preprint*, 2017. arXiv:1712.00547.
- [25] S. Rodriguez, J. Schaffer, J. O’Donovan, and T. Höllerer. Knowledge complacency and decision support systems. In *IEEE International Inter-Disciplinary Conference on Cognitive Methods in Situation Awareness and Decision Support (CogSIMA)*, 2019.
- [26] A. Shaban-Nejad. Semantic analytics for global health surveillance, April 2019. Retrieved on June 12, 2019 from <http://bit.ly/2ZkCzoT>.
- [27] J. Sowa. Representing and reasoning about contexts using IKL, September 2018. Retrieved on May 15, 2018 from <http://bit.ly/2xrtkIa> and <http://bit.ly/2y4h41v>.
- [28] S. Srihari. Explainable Artificial Intelligence: The Probabilistic Approach, April 2019. Retrieved on June 12, 2019 from <http://bit.ly/2UlmGES>.
- [29] R.D. Sriram and R. Sharma. Role of ontologies in explaining reasoning: An overview of explainable AI, November 2018. Retrieved on April 28, 2019 from <http://bit.ly/2TWYzUr>.
- [30] N. Tandon. Commonsense for Deep Learning, January 2019. Retrieved on June 12, 2019 from <http://bit.ly/2XE0cWS>.
- [31] I. Tiddi. Building intelligent systems (that can explain), March 2019. Retrieved on June 12, 2019 from <http://bit.ly/2SVUD41>.

- [32] I. Tididi, M. d'Aquin, and E. Motta. An ontology design pattern to define explanations. In *Proceedings of the 8th International Conference on Knowledge Capture*. ACM, October 2015. Article no. 3.
- [33] E. Topol. In *Deep Medicine: How Artificial Intelligence Can Make Healthcare Human Again*. Hatchett Book Group, New York, 2019.
- [34] E. Topol. High-performance medicine: The convergence of human and artificial intelligence. *Nature Medicine*, 25:44–56, January 2019.
- [35] A. Turano. Review and recommendations from past experience with medical explanation systems, February 2019. Retrieved on April 28, 2019 from <http://bit.ly/2WZfM0P>.
- [36] M. Underwood. Explanation use cases from retail credit card finance regulatory and service quality drivers, February 2019. Retrieved on April 28, 2019 from <http://bit.ly/2WLidE7>.

A Full Web Links

http://bit.ly/2ITKwFD	https://www.darpa.mil/attachments/XAIProgramUpdate.pdf
http://bit.ly/2KL6LVz	https://www.privacy-regulation.eu/en/r71.htm
http://bit.ly/2R9iD30	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Commonsense/Explanation--BenjaminGroszof_20190123.pdf
http://bit.ly/2RIViFQ	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Financial/Finance+Track+Session+1+--MikeBennett.pptx
http://bit.ly/2RiQSFz	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Commonsense/summit-physical-turing.pdf
http://bit.ly/2RqJQJ5	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Narrative/ProofAsExplanationAndNarrative--KenBaclawski_20190130.pdf
http://bit.ly/2SVUD41	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Narrative/BuildingIntelligentSystemsThatCanExplain--IlariaTiddi_20190313.pdf
http://bit.ly/2Scjvo6	https://s3.amazonaws.com/ontologforum/OntologySummit2019/XAI/ExplainableAI--WilliamClancey_20190220.pdf
http://bit.ly/2TM2141	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Overview/Summit-Overview-Session-2-DerekDoran_20181128.pdf
http://bit.ly/2TWYzUr	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Overview/Explainable-AI-and-Ontology--RaviSharma-RamDSriram_20181128.pptx
http://bit.ly/2Uhbxbwh	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Overview/IntroductionToCommonsenseAndExplanations--GaryBergCross-TorstenHahmann_20181205.pdf
http://bit.ly/2UlmGES	https://s3.amazonaws.com/ontologforum/OntologySummit2019/XAI/XAI--SargurSrihari_20190410.pdf
http://bit.ly/2VRt6DV	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Introduction/IntroductionToTracks--KenBaclawski_20190116.pdf
http://bit.ly/2VZa5zg	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Introduction/Commonsense-LaunchSession--GaryBergCross-TorstenHahmann_20180116.pdf
http://bit.ly/2WLidE7	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Financial/Explanation+Use+Cases+from+Retail+Credit+Card+Finance.pdf
http://bit.ly/2WTduUB	http://ontologforum.org/index.php/OntologySummit2019#Structure_and_Discourse
http://bit.ly/2WZfMOP	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Medical/ExplainableMedicalAI--AugieTurano_20190213.pptx
http://bit.ly/2XE0cWS	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Commonsense/CommonsenseForDeepLearning--NiketTandon_20190306.pdf
http://bit.ly/2YuPUe3	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Medical/MedicalExplanation--UgurKursuncu-ManasGaur_20190327.mp4
http://bit.ly/2ZkCzoT	https://s3.amazonaws.com/ontologforum/OntologySummit2019/Medical/SemanticAnalyticsForGlobalHealthSurveillance--ArashShabanNejad_20190417.pptx
http://bit.ly/2s9d4pH	https://www.darpa.mil/program/explainable-artificial-intelligence
http://bit.ly/2xrtkIa	http://ontologforum.org/index.php/ConferenceCall_2017_09_20
http://bit.ly/2y4h41v	http://ontologforum.org/index.php/ConferenceCall_2017_09_27
http://bit.ly/318aUbX	https://www.ihmc.us/users/acanas/Publications/ParticipatoryExplanation/ParticipatoryExplanation.htm