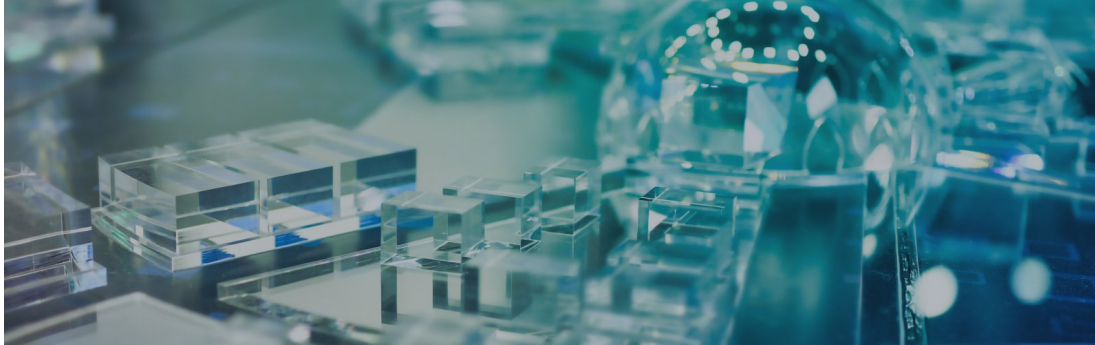


The Practical Problems of Explaining AI Black Box Systems

Originally published at <https://www.iflexion.com/blog/black-box-ai>

Published: May 28, 2020 | By Martin Anderson



If your dog kills someone, it's probably not the dog that will be subject to a lawsuit. Likewise, if you empower a brilliant but unpredictable thinking system with the ability to implement critical decisions, you will have to take responsibility for its actions.

Under this basic ethical template, the issue of accountability in machine learning systems has reached a critical junction: their cost-saving potential for the public and private sector increasingly conflicts with their opacity, either exposing governments and companies to prosecution and sanctions, or else demanding a more circumspect approach that could slow the evolution of AI to the same crawl that led to the two AI winters over the last fifty years that we discussed in [our machine learning overview](#).

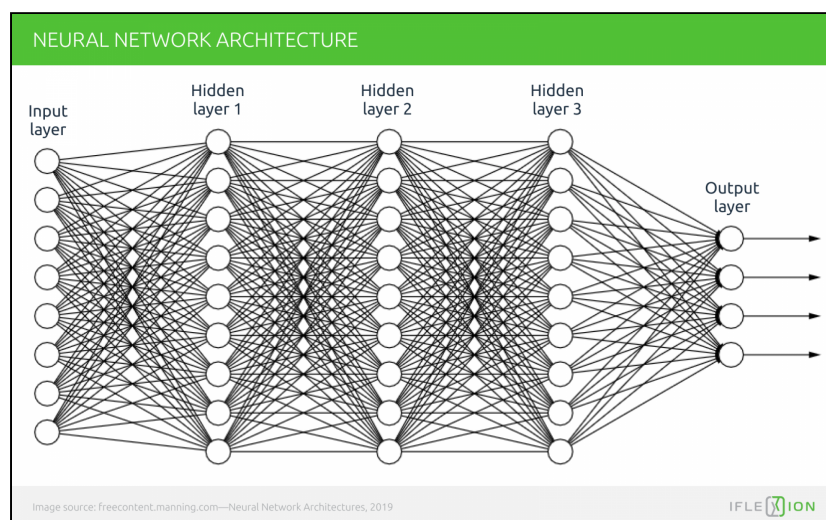
Where are we now with explaining the black box of AI? This article looks at both the state-sponsored and academic steps taken.

The Challenge of Deconstructing AI Systems

In response to this tension, the growing field of AI deconstruction is now seeking to explain the reasoning behind the decisions and outcomes of machine learning systems, and is gaining a strong commercial and political remit.

However, the labyrinthine, non-linear processes that machine learning operates on can be difficult to monitor in a meaningful way, and even more difficult to infer from outcomes.

The internal logic flow of a neural network is threaded through a complex web of simulated neurons that evolve and change in nature as they pass through interwoven layers.



Whether it's an image being splintered into pixels or a phrase sliced into words and character frequencies, data in a neural network is broken down into particulate instances and reconstructed systematically in the search for meaningful patterns. No mechanical observation method can easily keep track of the input information in this labyrinth, because the data itself is dissolved into a kind of 'digital gas' in the process.

Coming to Terms with Black Box AI

AI deconstruction must account for the social and practical integrity of the data. The prerequisite of [AI software development](#) is that an implicit, hidden or unconscious agenda in the acquisition and design of data, models, or frameworks can have a huge effect on the conclusions that AI draws from a dataset.

Furthermore, there are semantic disputes regarding a useful and meaningful definition of 'explainability' as it relates to AI black box systems. Oftentimes, the explanations are either so dumbed-down as to be inaccurate (or at least unhelpful), or else so technical that they are difficult to understand.

In the News: Negative Indicators for Empowering AI Systems

Empowered AI can do harm, and sometimes it does. But for the most part, the well-publicized failings of machine learning serve as an inhibiting factor for uptake or for extending the autonomy of machine systems:

- 'System limitations' in Tesla's onboard AI-based autopilot system contributed to the death of a driver in 2018.
- The advent of COVID-19 has compromised many brittle AI prediction models, such as those used by Amazon, which were not expecting such huge changes in data, and which have required manual intervention.
- The American Civil Liberties Union conducted a test showing that Amazon's Rekognition FR system matched up the faces of members of the US Congress with 28 known criminals, with an unpleasant additional racial bias.
- In 2014, Brisha Borden, a black minor offender aged 18 charged with one-time petty theft was rated higher for re-offending potential than a 41-year-old white career criminal by a Florida crime prediction AI, leading to further scandals around bias in machine learning in the police sector.
- In 2016, Microsoft's experimental machine learning chatbot was hurriedly taken offline when it turned into a genocidal racist, having been easily trained by online trolls attempting to influence its personality.
- A Chinese tech businesswoman was shamed as a jaywalker in Ningbo by an automated face recognition system which put her picture on the side of a bus as a cautionary example to offenders, even though the jaywalker was someone else.
- IBM's Watson medical AI framework offered 'unsafe and incorrect' cancer treatment advice, leading to the cancellation of the project in question and, eventually, layoffs in the Watson team.

Governmental and popular support for explainable AI has grown in the last five years in proportion to increased deployment and to governments' statements of intent around incorporating [AI in the public sector](#).

Concern has been registered over the last ten years regarding the future of big data and AI in sectors such as HR, public data management, credit scoring, welfare benefits access, health insurance, law enforcement, and vehicle systems, among many others.

Problems in Deconstructing Proprietary or State-Sponsored AI Algorithms

In 2017 Ed Felton, a professor of computer science and public affairs at Princeton, brought attention to the intellectual property aspect of explainable AI, wherein companies or governments may preclude reverse engineering or exposure of their algorithms, either for purposes of security, continuing exclusivity of a valuable property, or else national security.

Indeed, in seeking a balance between openness and security, AI research is sometimes ushered back into IP silos and away from useful public and peer scrutiny.

One review by the Committee on Standards in Public Life in the UK [contends](#) that 'it may not be necessary or desirable to publish the source code for an AI system'. They go on stating that '*[the] continuous refinement of AI systems could also be a problem if the system is deployed in an environment where the user can alter its performance and does so maliciously*'. The committee also found that the UK government is not adequately open about the extent to which machine learning is involved in public sector decisions.

In 2019, the research body OpenAI ironically decided not to release the source code of the text-creation algorithm GPT-2, declaring its ability to mimic human writing patterns as 'too dangerous'.

The impetus to withhold machine learning source code is at odds with the drive toward open standards and peer review in collaboration between the private sector and government — a further obstacle to explainable AI.

Additionally, publishing algorithms obtained through machine learning processes limits the field of AI deconstruction to a forensic, *post facto* approach, since an algorithm is only the outcome of a machine learning process rather than the architecture of the neural network itself.

Legal Quandary: An 'Act of AI'

The ethics and legality of reverse engineering, though always contentious^[1], were clearer before the age of machine learning: algorithms devised by humans were not only considered creative works that could enjoy IP protection, but lacked any concept of 'safe harbor': if the results were harmful, the creators were liable to assume responsibility and intent. Under this model, due diligence was relatively straightforward.

The case for AI-generated algorithms is less clear. Proving 'intent' is problematic when the creators have an imperfect understanding of how their own algorithm was formulated. Thus, the algorithm distributors remain *responsible* for its output, without understanding exactly how the negative results transpired or were formulated by the machine learning process — a truly toxic position.

Where AI causes an unintended and negative outcome, there is a legal vacuum with regard to 'intent', since the creators presumably did not wish to cause harm. Yet 'neglect' is not applicable in the absence of any available system or process that could have prevented it.

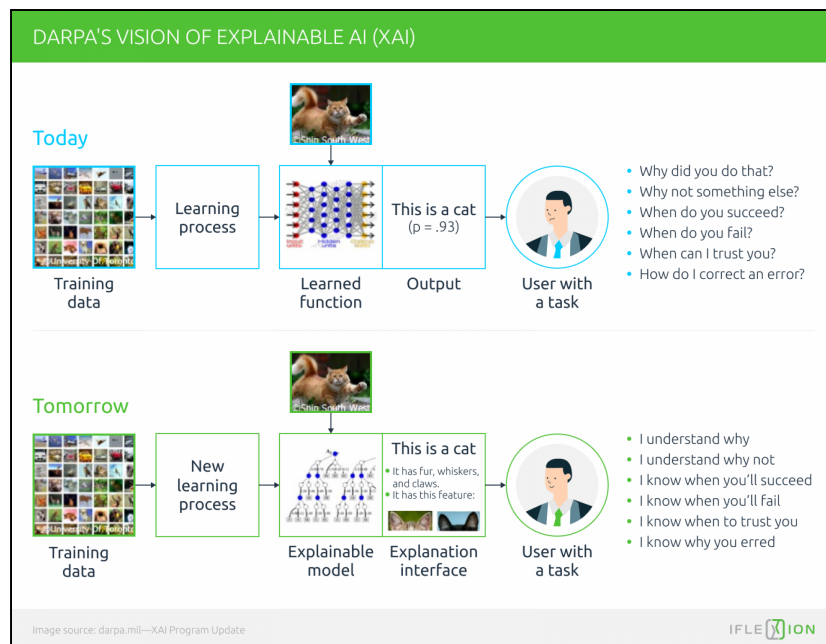
Without definitive legislation, such events could be interpreted as 'acts of AI' in the same sense that insurance companies define an 'act of God' as an uninsurable risk.

Thus powerful interests are both enabled and threatened by the possible outcomes of their own opaque AI systems, swept along by the economic impetus for automation and data analysis but exposed by the lack of control mechanisms and reproducible methodologies.

Without definitive legislation, there could be 'acts of AI' in the same sense that insurance companies define an 'act of God' as an uninsurable risk.

DARPA's Explainable AI (XAI) Program

Though various pieces of US legislation address the issue of AI accountability to some degree, and though there is a notable ad hoc body of state-sponsored academic research in the field, concerns about black box AI have manifested primarily in DARPA's Explainable Artificial Intelligence (XAI) initiative in the USA.



Founded in 2017, the XAI project has struggled to find effective or applicable remedies for the opacity and inscrutability of machine learning reasoning processes. Over the course of three years, its emphasis has switched from short-term architecture intervention to a collaborative academic effort to quantify, name and rationalize core concepts and possible solutions, with an emphasis on Explainable AI Planning (XAIP).

The greater body of DARPA's literature centers on the forensic deconstruction of AI model decisions, with inferential analysis of AI outcomes emphasized over in-model tracking techniques.

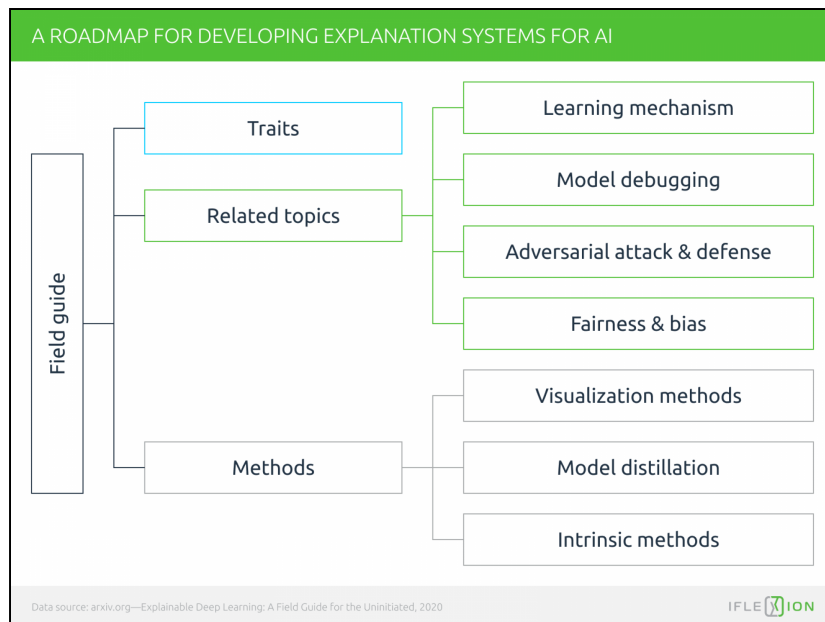
The XAI initiative proceeded from the assumption (or at least the hope) that automated reporting could be integrated into machine learning processes while keeping them performant. But [the 2020 survey of DARPA's curriculum](#) by IBM and Arizona State University reflects that the project's research into automating AI explainability has evolved into a more taxing examination of potential robot/human collaboration models, with automated analysis perceived as 'computationally prohibitive'.

General Academic Research into Explainable AI

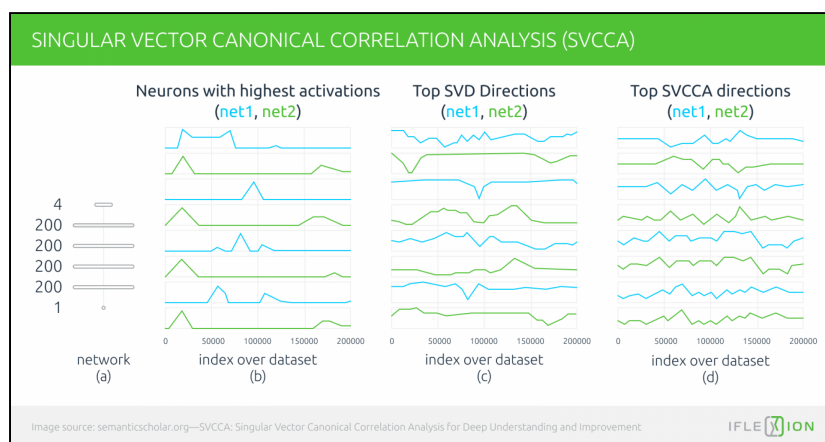
In April 2020, researchers in the US and the Netherlands released [a comprehensive overview](#) of current initiatives in explainable deep learning. The report reveals how metaphysical the field of AI deconstruction currently is, and that the forward direction centers more on human-focused planning strategies rather than 'static analysis' techniques applied against an existing neural network.

The researchers break down the challenge into three core areas that cover a great deal of interrelated research:

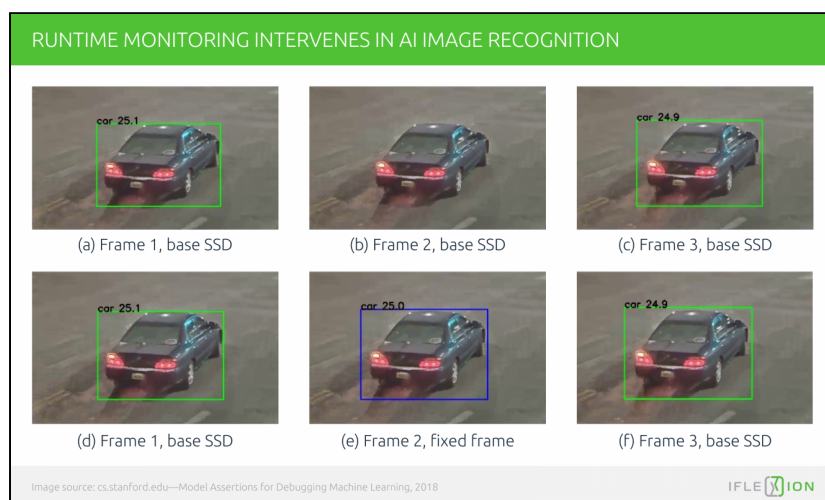
- **Traits**, wherein the objectives of AI explanations are examined, and terms for explainability are defined.
- **Related Topics**, where XAI is stress-tested against analogous research fields.
- **Methods**, which examines and questions recent academic assumptions in the search for foundational principles for XAI.



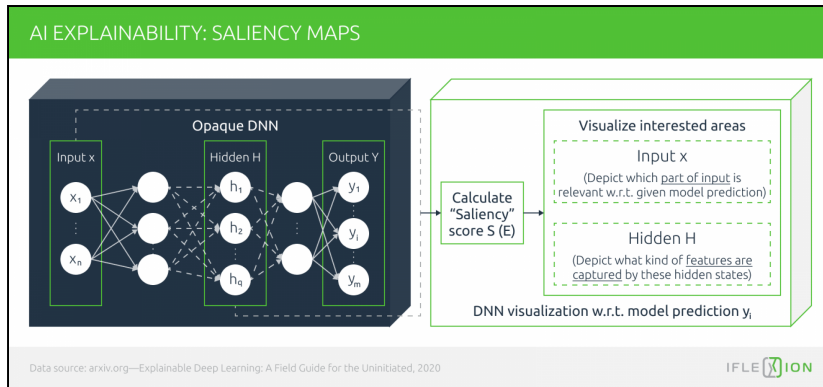
One approach noted in the research is Singular Vector Canonical Correlation Analysis (SVCCA), wherein individual neurons in a DNN are mapped into a vector set that can reveal layer activations at runtime.



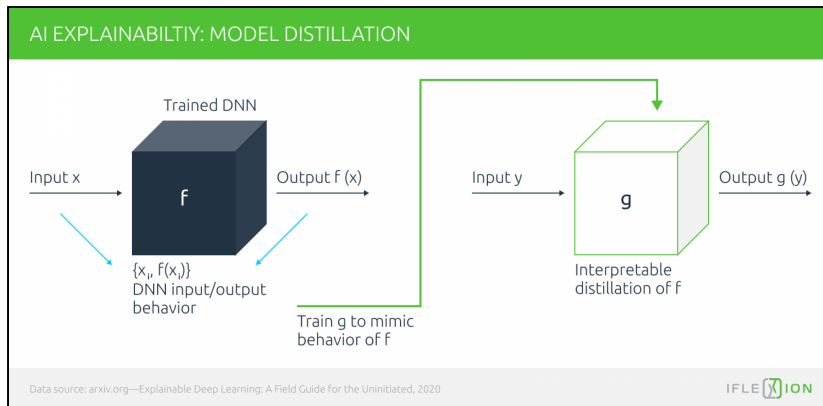
Though the literature on real-time model reporting is less extensive, research out of the Stanford DAWN Project proposes 'Model Assertions' as a run-time method of monitoring and intervening in the live processes of a deep neural network. However, the technique's success with image processing pipelines may be difficult to replicate in more abstract realms such as language processing or [sentiment analysis](#).



Visualization methods can also help map model features that cause high stimulation in a neural network, and are facilitated either by back-propagation or perturbation-based visualization.



Model Distillation is a post-facto technique wherein the output of a DNN is used as input data for a secondary, explicatory DNN.



A number of 'intrinsic' methods suggest possible ways that deep learning models can self-report on the logic of their process pipelines:

- **Attention Mechanisms** can assign tracking capabilities to specific inputs or else force the DNN to ask questions before it continues processing data, and current methods include *Single-Modal Weighting* and *Multi-Modal Interaction*.
- **Joint Training** proposes several methods of 'chaperoning' the data with functions designed to account for what is happening to it in a neural network. Current methods include *Text Explanation*, where a DNN is augmented with an explanation generation component; *Explanation Association*, where data is associated with human-interpretable concepts and objects, in the hope that the evolution of this secondary layer makes sense of the primary layer; and *Model Prototype*, designed for classification models, where case-based reasoning forms an association between input data and prototypes observations in the dataset.

Conclusion

In one famous example of sci-fi satire, the fictional supercomputer Deep Thought took 7.5 million years to compute the answer to 'the meaning of life', which turned out, unhelpfully, to be '42'.

According to the latest literature, life is imitating art, as the locus of research into XAI turns toward asking better questions, defining clearer terms, and conceiving new AI frameworks whose reasoning can be made more explicable than that of current machine learning approaches.

In the balance, the progress in interpretive post-facto XAI methods and principle development remains far ahead of in-process explanation mechanisms for deep learning.