

The Problems of Policing Artificial Intelligence

Originally published June 29, 2017 at <https://www.aiaa.net/artificial-intelligence/articles/the-problems-of-policing-artificial-intelligence>

Wrong assumptions by both machines and administrators are slowing progress of AI-driven crime prediction



Photo by [Matt Popovich](#) on [Unsplash](#)

A new study around transparency in civic machine learning solutions reveals some of the political and conceptual issues that seem likely to grab headlines in years to come.

The research, undertaken by Michael Veale of the Department of Science, Technology, Engineering & Public Policy (STeEP) at University College London, provides excerpts of interviews with five anonymous government contributors and ex-contributors over three continents.

The 'pre-crime' machine learning frameworks discussed are used for projecting areas and individuals likely to become the focus of police intervention in the near future, based on historical data.

One anonymous analytics leader featured found that he had to take ethical advice on the use of various in-house machine learning systems around individual offender/victim prediction and geospatial crime prediction:

"[We] had guidance from the ethics committee on this. We were to work down the list, allocating resources in that order, and that's the way they told us would be the most ethical way to use them... It's also important to make clear that the professional judgement always overrides the system. It is just another tool that they can use to help them come to decisions."

A software modeller at one of the AI projects discussed the problems of scalability at the point where promising AI-driven intelligence-gathering systems might be rolled out to larger police resources, observing that the transition from nursery to policy often gets handled 'by someone not qualified enough' for the task.

Further problems emerge when attempting to scale up or roll out predictive AI systems which seem to have proven effective in the field. An in-house software engineer in one police district, who built various models to predict human trafficking 'hotspots', recounts an instance of failure-to-scale:

"Thankfully we barely have any reports of human trafficking. But someone at intel got a tip-off and looked into cases at car washes, because we hadn't really investigated those much."

"But now when we try to model human trafficking we only see human trafficking being predicted at car washes, which suddenly seem very high risk. So because of increased intel we've essentially produced models that tell us where car washes are."

"This kind of loop is hard to explain to those higher up."

When analytics users in such systems have an unclear understanding of how one of the models involved works, they can easily become 'wary of picking up protected characteristics,' according to one interviewee, indicating further need for granular policies around the machine systems. But since so many ML systems are prototype in nature, or at least highly customized or non-standardized, even more general policy can be hard to obtain and ratify.

Other interviewees warn of the problems of [automation bias](#), where an often-mistaken assumption can arise that the worker is serving the machine's judgement rather than evaluating it and using it, quite correctly, as questionable intelligence.

The paper recounts how regional intelligence officers — who correlate local news, reports, and other information sources into intelligence strategies — were asked to compare their own judgement to an AI system's estimation of crime probability:

"They might say they know something about the offender for a string of burglaries, or that building is no longer at such high risk of burglary because local government just arranged all the locks to be changed. [...] We also have weekly meeting with all the officers, leadership, management, patrol and so on, with the intelligence officers at the core. There, he or she presents what they think is going on, and what should be done about it."

But the use of such systems brings police authorities into very sensitive areas in terms of PR, raising ethical issues around using such AI-generated lists as policy guidelines.

In terms of transparency, policy makers need to straddle a line between public accountability and the risk of exposing a valuable AI-driven intelligence system to external gaming. One lead at a government analytics department concerned with investigating tax fraud expressed fears that the usefulness of the information might diminish if model weights were made public as part of more general transparency policies.

The report gives the sense that civic police departments don't appear to have the kind of leverage necessary to hide essential facets of their Machine Learning strategies behind national security policies such as the Patriot Act (though the USA was not one of the countries involved in the [research](#)).