

Los límites estadísticos de la vigilancia masiva no sirven para detectar individuos, sino para criminalizar colectivos

Por: Javier Sánchez Monedero. Rebelión. 27/11/2019

Después de los atentados del 11-S de 2001, Estados Unidos inauguró la senda del espionaje masivo de la ciudadanía en busca de potenciales terroristas. La premisa, mantenida e impulsada como mantra gracias al avance del big data, era y es: cuantos más datos tengamos en número y variedad más probabilidades tendremos de detener a los terroristas antes de que actúen. Desde entonces, en el debate respecto a herramientas de vigilancia y estrategias preventivas de seguridad se introdujo una salvaguarda narrativa en forma de dilema inevitable: seguridad o libertad. Y ese dilema ha servido para calmar las resistencias a una agenda política que normaliza el control y la vigilancia, además de criminalizar a grupos sociales.

La oposición a estos despliegues se plantea a menudo desde los derechos sociales y humanos, pero rara vez incorporamos al debate público un análisis estadístico que nos demuestra que es imposible que la mayoría de estas propuestas funcionen para el que se supone que es su propósito. Es decir, no hay ningún dilema seguridad-libertad que resolver.

Desde hace años, la Unión Europea cuenta con varios sistemas tecnológicos para el control de [personas refugiadas y migrantes](#) mientras los estados miembro y la propia UE no paran de financiar o implementar planes piloto de vigilancia de sus ciudadanos. Por ejemplo, Gales puso en marcha en 2017 un sistema de reconocimiento facial para evitar el acceso a potenciales criminales a eventos deportivos.

El objetivo era detectar personas que pudieran suponer un peligro entre las asistentes a estos eventos. Sin embargo, cuando el sistema se desplegó en varios eventos deportivos, como la final de la Champions League, su rendimiento fue mucho menor que en las pruebas de laboratorio: [un 92% de las 2,297 personas detectadas como criminales, y a las que se les impidió acceder al espacio](#), no tenían antecedentes. El problema de los falsos positivos con el reconocimiento facial también dio que hablar en EEUU cuando [el sistema reconocimiento de Amazon identificó a 28 congresistas como delincuentes](#).

En las mismas fechas, un equipo de investigación de la Universidad de Granada (UGR) [presentaba un detector de mentiras basado en la monitorización de la temperatura de la nariz](#). Entre los ejemplos de aplicación los inventores proponían su uso en campos de refugiados para saber “*cuál es el objetivo real de las personas que tratan de cruzar las fronteras entre países*”. El artículo asegura alcanzar una tasa de detección de mentirosos del 85%.

Obviando la [unanimidad científica](#) en que no existe ninguna base teórica o experimental para pensar que sería posible construir un detector de mentiras y que tampoco podemos diseñar un experimento sólido de validación (todos los trabajos que proponen detectores de mentiras estudian su rendimiento con personas que interpretan un papel), estas noticias son un ejemplo perfecto del mencionado falso dilema. Bajo este, se reducen las resistencias a someter a un grupo social a una prueba claramente criminalizadora con la promesa de detectar terroristas de forma eficaz y avalada por instituciones académicas.

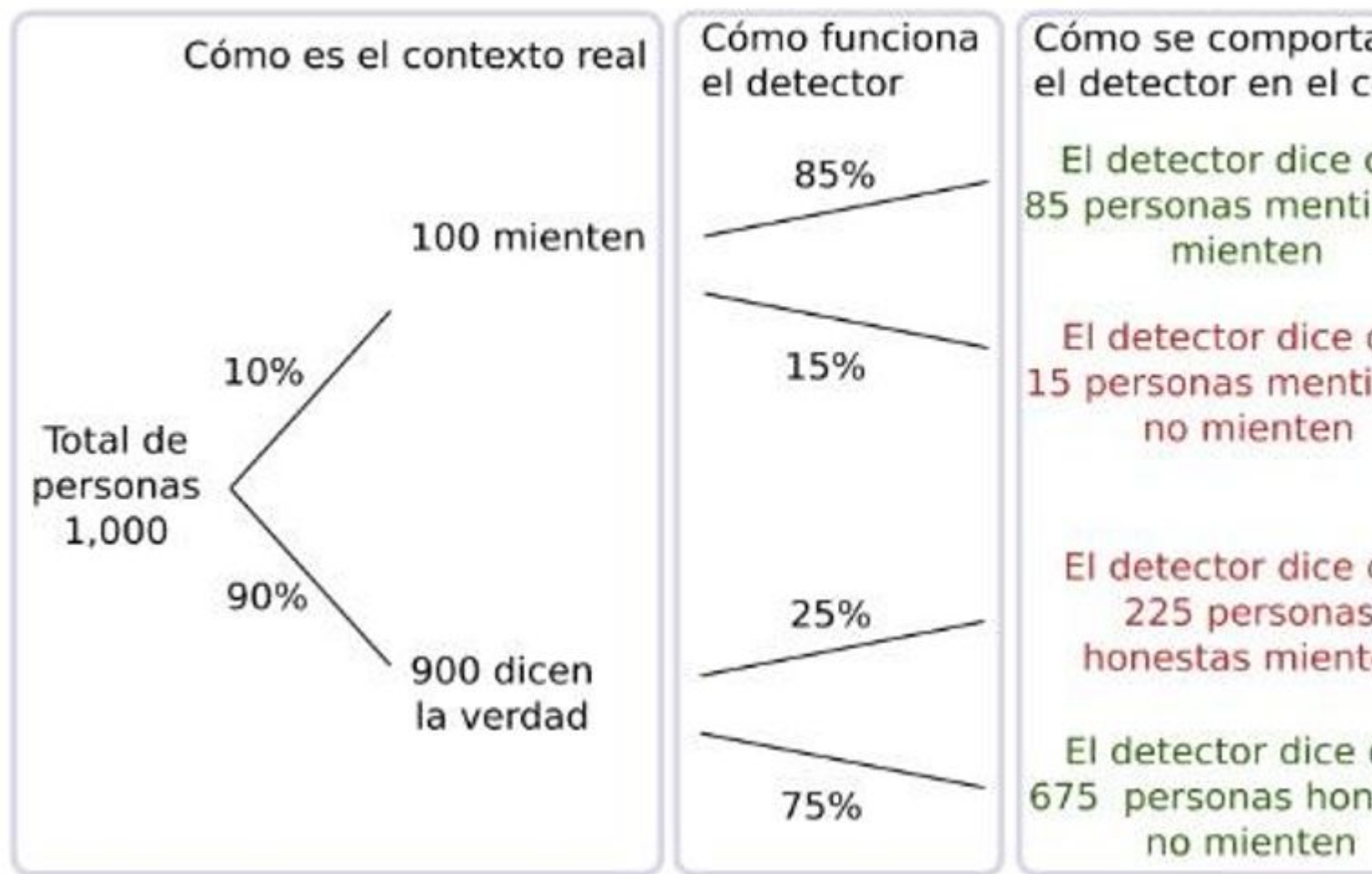
Pero incluso en el caso de que estos detectores de mentiras funcionasen y diéramos por buenas sus pruebas de laboratorio, así como la hipótesis de que en los campos de refugiados se esconden terroristas, ¿qué pasaría si pudiéramos demostrar que es altamente improbable, por no decir imposible, que muchas de estas propuestas funcionen? La estadística nos da herramientas para ello.

Volviendo al detector basado en el *Efecto Pinocho*, ¿tiene sentido la propuesta del investigador para detectar personas que se hagan pasar por refugiadas? Para esto necesitamos primero entender los resultados de laboratorio y después contextualizarlos en el problema real.

El experimento

Antes de comenzar, es importante señalar que el hecho de que una persona mienta es un evento, mientras que el hecho de que el detector diga que miente es otro evento separado. El experimento de laboratorio, con un grupo de estudiantes que interpretaban unos roles, arrojó unos resultados de 85% de detección de mentirosos y 75% de detección de personas que dicen la verdad. Esto significa que el 15% de las personas que mienten no son detectadas y que el 25% de las que dicen la verdad son clasificadas como mentirosas. A simple vista parecen unos resultados prometedores, pero sería interesante saber si en la práctica real los resultados de rendimiento se mantendrían.

La estadística nos da herramientas que nos sirven para contextualizar el comportamiento de una prueba como el detector de mentiras teniendo en cuenta cómo de frecuente es el evento que queremos detectar. Por ejemplo, si tenemos 1.000 personas de las cuales 100 mienten, ¿qué hará nuestro detector de mentiras? Si hacemos las cuentas, 85 de las personas mentirosas serán detectadas y 15 no, y de las 900 personas honestas el detector dirá que 675 dijeron la verdad y que 225 mentían. En este punto, cabe destacar que el detector habrá clasificado 310 personas (85 + 225) como mentirosas, y que mirando sólo el detector de mentiras no podemos saber nada más, de modo que habría que investigar a estas personas para encontrar a las que de verdad mienten, que se espera que sean un 27% de todas las clasificadas como mentirosas. Formalmente a esta corrección sobre el rendimiento se le conoce como *valor predictivo positivo* o *probabilidad posterior* y es imprescindible para valorar utilizar una prueba en un contexto real.



Podemos repetir este análisis para saber si el detector de mentiras de la Universidad de Granada podría encontrar terroristas haciéndose pasar por refugiados. Para ello utilizaremos los datos del *think tank* de extrema derecha Heritage, que asegura que 44 demandantes de asilo, de entre 4 millones, estuvieron relacionados con algún evento de terrorismo en Europa entre 2014 y 2017, esto es el 0.0001% del total. Repitiendo los cálculos anteriores, 978.886 personas serían clasificadas como mentirosas, entre las cuales estarían 36 terroristas y el resto serían personas inocentes. Es decir, habría que investigar a casi 1 millón de personas en extrema vulnerabilidad para encontrar a esos 36 terroristas. Si calculamos el valor predictivo positivo este nos da un 0,0037%, una cifra que impresiona menos que el 85% de precisión anunciada por la UGR.

Es importante destacar que en nuestro análisis hemos hecho una serie de asunciones que representan el mejor de los casos para la propuesta de la UGR: los

detectores de mentiras existen, los experimentos de laboratorio se parecen al entorno real de pruebas y por tanto el rendimiento se mantiene, hemos sobreestimado la cifra de refugiados implicados en terrorismo y hemos dado por hecho que los terroristas no adoptarían ninguna medida para evitar ser detectados.

Después de nuestro análisis queda claro que la propuesta es inviable, no sólo para la búsqueda de terroristas sino para la búsqueda de muchos otros eventos, ya que para alcanzar un valor predictivo positivo de más del 50% el evento que buscamos tendría que cumplirse en al menos el 25% de la población. Sin embargo, cualquiera que haya leído la nota de prensa de la UGR habrá concluido que la entidad propone, con el aval de la (pseudo)ciencia, que hacer pasar por una prueba criminalizadora a los refugiados es una buena idea.

La historia de los detectores de mentiras está plagada de casos fallidos lo cual no ha impedido que, recurrentemente, países de la UE o la propia Comisión Europea financien programas piloto, de los que, casualidad, casi nunca encontramos información sobre los resultados prácticos. Esto ha sucedido con la implantación de AVATAR en EEUU y Hungría para entrevistar a migrantes, el proyecto de la Universidad de Bradford y QinetiQ en 2011 para buscar defraudadores en las ayudas sociales, y más recientemente con iBorderCtrl, que, financiado con 4,4 millones de euros por la UE promete agilizar el paso fronterizo sometiendo a personas sin un visado a un detector de mentiras supuestamente basado en inteligencia artificial.

Ahora que sabemos que estos sistemas no son fiables, y demostrado estadísticamente que no pueden funcionar en la búsqueda de eventos poco frecuentes, cabe preguntarnos qué lógicas persigue la UE con estos pilotos: ¿responden a una creciente criminalización de las migraciones? ¿Son una herramienta tecnoburocrática que pretende deshacerse de migrantes bajo el supuesto aval de objetividad que otorga la ciencia? ¿Tienen que ver con una industria de la seguridad parasitaria de Bruselas? Por lo pronto, European Dynamics, la empresa detrás de iBorderCtrl y muchos otros proyectos financiados por la UE, sigue sin publicar las memorias completas de los proyectos, incluyendo los documentos del comité asesor de ética, y ya tiene en marcha otro proyecto europeo que promete detectar comportamientos de riesgo en migrantes que cruzan las fronteras. Sería interesante saber por qué la Comisión Europea ha calificado como “caso de éxito” el detector de iBorderCtrl cuando sabemos que se basa en pseudociencia y sería muy difícil que aportase alguna utilidad incluso con todas las

hipótesis en su favor. Sería oportuno analizar por qué hay consultorías dedicadas en exclusiva a escribir proyectos de investigación y también sería interesante seguir el rastro del dinero, pero eso sería otro tipo de artículo.

[LEER EL ARTÍCULO ORIGINAL PULSANDO AQUÍ.](#)

Fotografía: Rebelión

Fecha de creación

2019/11/26