

> linguistiques pour évaluer la crédibilité, par rapport aux indicateurs paraverbaux (le ton de la voix, la longueur des pauses, etc.) et aux comportements non verbaux (les gestes, les postures, etc.) : grâce à « une métaanalyse » portant sur ces recherches, Bella DePaulo, de l'université de Virginie, a par exemple montré en 2003 que les récits de menteurs sont moins détaillés que les versions d'individus sincères. En d'autres termes, si l'on cherche à tirer une histoire au clair, mieux vaut bien écouter et se concentrer sur ce qui est dit plutôt que de chercher à observer les indices corporels ou vocaux. Là encore, *iBorderCtrl* gagnerait probablement à creuser dans cette direction, en incluant des indicateurs linguistiques.

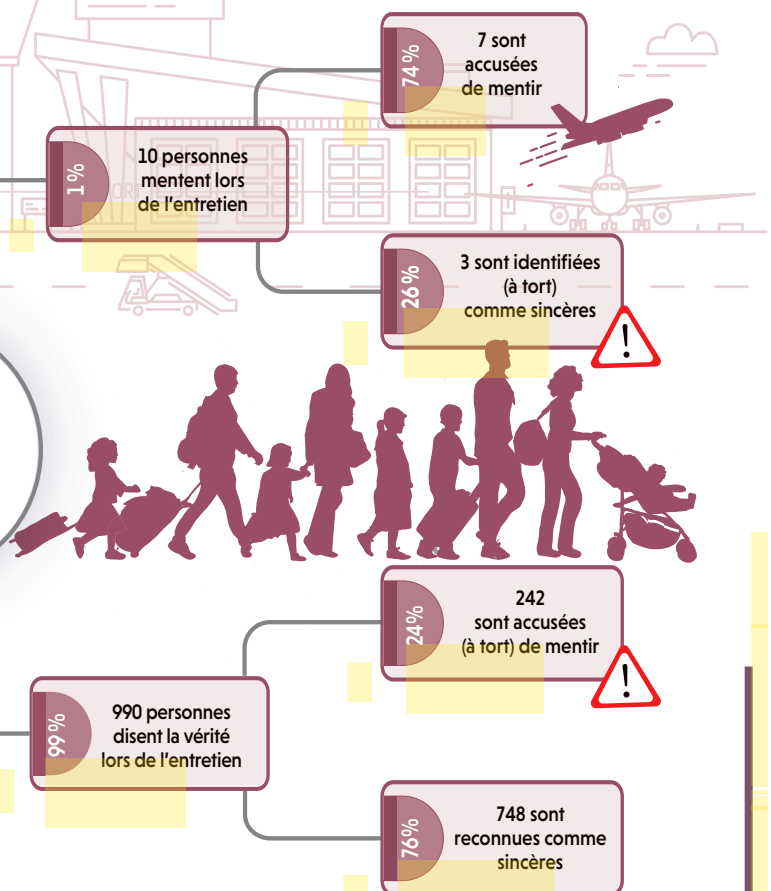
Un autre problème est que ce système a été entraîné et préalablement testé sur un nombre assez restreint de personnes, environ une trentaine. Cela peut se traduire par une limitation bien connue en intelligence artificielle, nommée l'« *overfitting* ». C'est-à-dire que le modèle risque de ne pas être généralisable à d'autres individus, notamment ceux qui parleraient une autre langue ou s'exprimeraient différemment.

D'après le site officiel de *iBorderCtrl*, l'outil a tout de même été éprouvé dans des conditions quasi réelles en 2019, sur la base du volontariat, aux frontières de la Hongrie, de la Grèce et de la Lituanie. Ses développeurs n'ont communiqué aucune donnée sur les résultats,

ACCUSÉS À TORT DE MENTIR

On évalue l'efficacité d'un outil de détection automatique des mensonges avec deux taux : la proportion des menteurs qui sont identifiés comme tels et la proportion de personnes disant la vérité qui sont effectivement reconnues comme sincères. Pour *iBorderCtrl*, ces taux seraient respectivement de 74 % et 76 %, selon des tests préliminaires. Comme on le voit sur la figure, pour 1 000 voyageurs qui répondraient à une question, 3 menteurs passeraient à travers les mailles des filets (en supposant qu'il y a 1 % de menteurs), mais 242 personnes sincères seraient accusées à tort. Ce nombre élevé s'explique par le fait que dans une application comme le contrôle des frontières, la grande majorité des gens (99 % dans le modèle proposé ici) disent la vérité, de sorte que même un taux de 76 % conduit à de nombreux « faux positifs » – ce taux signifie que lorsque vous êtes sincère, vous avez environ une chance sur quatre d'être accusé de mentir. Alors, imaginez un instant ce que cela donnerait avec les centaines de millions de personnes qui franchissent chaque année les frontières de l'Europe, et qui devraient répondre à de nombreuses questions !

1 000
voyageurs



mais la mésaventure de la journaliste italienne Ludovica Jona donne une idée de ce qui attendrait nombre de voyageurs avec un tel détecteur automatique de mensonges. Elle a réussi à expérimenter le système lors de cette phase de tests. Pour un résultat édifiant: le système l'accusait de mentir sur sa date de naissance, sa citoyenneté, le lieu de délivrance de son passeport et sa destination! Elle a finalement été considérée comme «à risque», alors qu'elle avait été sincère lors du préenregistrement.

DES QUESTIONS ÉTHIQUES ET SOCIÉTALES

Car n'oublions pas que 70% de réussite, cela veut dire 30% d'échec, ce qui est loin d'être négligeable. Concrètement, ces erreurs consistent à considérer soit un menteur comme sincère, soit une personne sincère comme menteuse. Et même si l'on parvenait à améliorer considérablement la fiabilité, cela signifierait que sur les centaines de millions de voyageurs qui franchissent chaque année les frontières de l'Union européenne, plusieurs millions seraient soupçonnés à tort de mentir! Tandis que bien des personnes pouvant présenter un risque sécuritaire passeraient à travers les mailles du filet...

Ce n'est qu'une des nombreuses questions éthiques que soulève l'utilisation de l'intelligence artificielle dans ce contexte. Quelles seraient les conséquences d'une telle gouvernance technologique? Comment la concilier avec les libertés fondamentales et le respect de la vie privée? En développant de tels outils, n'ouvre-t-on pas la porte à une surveillance de masse d'une ampleur jamais observée?

Pour l'heure, *iBorderCtrl* n'en est qu'au stade de la recherche et on ignore quand il sera mis en application – ni même s'il le sera vraiment. Mais une tendance de fond se dégage quant à l'accélération de la mise en place de dispositifs de surveillance dans le monde. La Chine dispose déjà de 200 millions de caméras de surveillance et compte atteindre le cap des 600 millions dans un an. Des algorithmes sont déjà utilisés pour détecter les baisses de motivation des ouvriers dans les usines ou des élèves à l'école. En France, la reconnaissance faciale va être testée dans les lieux publics pour une durée de un an, censément sous le contrôle de chercheurs et de membres de la société civile.

Cette situation appelle deux remarques essentielles. Premièrement, au fil de ces développements, il sera capital de rester au plus près des données scientifiques concernant ces dispositifs, pour avoir conscience de ce qu'ils peuvent faire et de ce qu'ils ne peuvent pas faire. Cet article poursuit cet objectif en montrant clairement qu'aucun système de ce type ne pourrait décemment être mis en application à l'heure actuelle.

C'est une situation entièrement nouvelle pour notre constitution mentale

Deuxièmement, il s'agira de décider avec clairvoyance des domaines de la vie publique où l'interdiction de mentir sera promulguée et punie avec le renfort de moyens technologiques. Le cerveau et la psyché humaine ont évolué en partie grâce à la capacité de taire certaines pensées et d'énoncer des faits différents – clairement, de mentir (c'est notamment une part importante de ce qu'on appelle la «théorie de l'esprit»).

Cette situation est donc entièrement nouvelle pour notre constitution mentale et il ne faut pas sous-estimer ses risques potentiellement destructeurs. Si l'humain s'est accommodé depuis des millénaires de l'obligation morale de dire la vérité dans certaines circonstances (on dit bien aux enfants de ne pas mentir, après ils font ce qu'ils veulent...), rien ne nous garantit qu'il saura le faire avec l'obligation physique de dévoiler le fond de sa pensée, sans aucune échappatoire. Un homme transparent serait-il encore un homme? Il ne faudrait pas que les recherches en psychologie autour de ces questions prennent trop de retard sur celles qui visent la seule détection du mensonge. ■

BIBLIOGRAPHIE

J. Sánchez-Monedero et L. Dencik, **The politics of deceptive borders : « Biomarkers of deceit » and the case of iBorderCtrl**, *Information, Communication & Society*, en ligne le 3 août 2020.

S. Jordan et al., **A test of the micro-expressions training tool : Does it improve lie detection ?**, *J. Investig. Psychol. Offender Profil.*, vol. 16, pp. 222-235, 2019.

B. Elissalde et al., **Le mensonge - Psychologie, applications et outils de détection**, Dunod, 2019.

S. Porter et al., **Secrets and lies : Involuntary leakage in deceptive facial expressions as a function of emotional intensity**, *Journal of Nonverbal Behavior*, vol. 36, pp. 23-37, 2012.

à lire



THEMA L'intelligence artificielle

Une sélection des meilleurs articles publiés dans *Pour la Science* ou *Cerveau & Psycho* sur l'intelligence artificielle, ses succès, ses limitations, ses risques.

Hors-série numérique en vente seulement sur : boutique.pourlascience.fr