

AI deep fake fraud exposed - five key lessons and how we need to respond

 12 April 2024  Paul Wainwright

Whether you believe the mainstream news channels, are informed via social media or retain genuine scepticism about the news agenda, like it or loathe it, fake news is an inherent part of our daily information diet. This has recently taken various forms involving celebrities and politicians being manipulated to look good or bad, using the latest AI technology. Reputations can be easily lost as a result.

However the recent case of AI deep fake fraud identified in an article published by CNN ([Deepfake Cfo scam](#)) has highlighted an even more sinister development in the use of AI technology to commit fraud where AI deep fake technology intrudes in ordinary business dealings.

The CNN article sets out the background of a scam where a genuine worker in finance was hoodwinked by a new level of deception. It has the hallmarks of CEO Fraud but with fake moving and interactive images rather than a shouty email from a false (or hacked) senior exec.

Although the details are relatively scant, for good reason, some key points can be drawn from the case. Credit too goes to the wider investigations carried out by the Hong Kong Police into such fraudulent enterprises which highlight the issues which we all face on remote video business interactions:

1. Fraudsters used deepfake technology to pose as the company's chief financial officer in a video conference call, tricking a finance worker at a multinational firm into paying out the equivalent of \$25 million. It should now be recognised that people in leadership positions with online profiles – and public recordings of them in podcasts or video settings are more exposed to the risk of ID theft which can be used with AI tools in voice and images for fraud.
2. The worker grew suspicious after he received a message that was purportedly from the company's UK-based chief financial officer. However, the worker attended the video call with what he thought were several other members of staff, but all of whom were in fact deepfake recreations, and this seems to have allayed his doubts and laid the platform for the fraud. Multiple deep fake AI identities had therefore been used to perpetrate this fraud. The complexity of the video conference suggested a high level of security penetration and information gathering from within the organisation.
3. In the wider investigation fraudsters used eight stolen Hong Kong identity cards to make 90 loan applications and 54 bank account registrations between July and September last year. On at least 20 occasions, AI deepfakes had been used to trick facial recognition programs by imitating the people pictured on the identity cards. Banks are also in line to be exploited and AML checks of themselves seem to lack to sophistication to stop false accounts being set up. Independent verification in real time for customer payments should be the norm especially where new account or supplier details are being adopted.
4. Training and due diligence is vital – don't trust what you read, hear or see even from senior management. The old adage of having suitable checks and division of responsibility, remain paramount. The case highlights the importance of verifying the identity of the person making the request, and the request itself being backed up with authenticating paperwork and verification documents confirming account details. If the worker had taken the time to verify the identity, and the chain of correspondence, and the payee, via the bank (especially given the amount) the fraud could have been prevented.
5. On perhaps a more controversial note, and one for further debate no doubt, the discovery of this AI deep fake fraud also highlights the need for the imposition of an ethical dimension to the use of AI technology. The increased availability of off the shelf technology to

replicate such scams strongly suggests that regulation of AI technology is potentially the only way to prevent its misuse. The EU's proposed AI Act (likely to be in form from around June 2024) would categorise such activity of AI systems as High Risk and likely to be banned. Allowing AI regulators to identify the level of risks associated with such applications would be a start.

The discussion on the use of AI technology and the potential risks associated with deep fake AI is relatively new. It must also factor in the complex arguments likely to come from the US and move beyond the "guns don't kill people, people do" rhetoric. The fact that social media giants are presently not considered the authors of the offensive and harmful material (under Section 230 of the Communications Decency Act 1996 in the US) will change in the UK with the regulations due to come following Royal Assent of the Online Safety Act. It remains to be seen how effective the provisions for proportionate measures to protect users from fraudulent material will be. Failure to govern harmful online activity on the basis of free speech or to curtail harm by acting in tandem with a comprehensive counter fraud strategy may be considered by some as a license to commit fraud.

Overall, the case and discovery of AI deep fake fraud of such magnitude serves as a wake-up call for many organisations to increase their security measures and take staff awareness of the potential risks of AI deep fake technology seriously. By learning from this case and implementing the necessary measures, organisations can better protect themselves from fraud and other security threats.

Key contact



Paul Wainwright

Partner

paul.wainwright@brownejacobson.com

+44 (0)121 237 4577

Related expertise

Business crime and fraud

Corporate

Criminal compliance and regulatory

Technology disputes