

Thema: Kansen en bedreigingen in AI en cybersecurity

- ◆ Crimineel gebruik van AI
- ◆ Samen beslissingen nemen met Explainable AI
- ◆ Column Kwantummechanica uitleggen aan je hond



ISOPlanner

Eenvoudig Compliance Management in **Microsoft 365**

Hoe Gemeente Waterland structuur en controle heeft bij de invoering van BIO-maatregelen

De gemeente had een duidelijk doel: controle houden op de uitvoering van maatregelen om te voldoen aan de Baseline Informatiebeveiliging Overheid (BIO).

Door de samenwerking met ISOPlanner:

- ✓ Heeft de gemeente een ISMS dat overzicht biedt in benodigde acties.
- ✓ Bespaart de gemeente tijd door gebruik te maken van voorbeelddocumenten.
- ✓ Borgt de Microsoft-integratie het doorvoeren van maatregelen.

ISOPlanner is de enige integrale compliance oplossing in Microsoft 365. Uiteraard ook voor ISO 27001, NEN 7510 en NIS2.



Snelle implementatie door meegeleverde templates en naadloze integratie met Teams, SharePoint, Outlook en Power Automate.

We kozen voor een snelle implementatie waarbij veel voorwerk was gedaan. En dat aansloot bij onze beveiligingseisen zoals single sign-on. Ook de Microsoft-integratie was een veriste.

Jimmy Voskuil
Gemeente Waterland



Kijk op **www.isoplanner.app** voor meer informatie en jouw gratis proefperiode.



Ik kan het niet meer horen!



Chris de Vries

Nee, dit gaat niet over de overvloed aan nieuwjaarswensen, maar over het begin van het artikel over crimineel gebruik van AI, namelijk: 'AI. Soms kan ik de term niet meer horen'. En juist daarover gaat dit nummer. De een ziet de enorme kansen en verwijst daarbij naar de 'Cyber Resilience Act (CRA)', terwijl de ander juist huiverig wordt, denkend aan ethische dilemma's en de vraag: 'Want wie kiezen we... en wie kunnen we nog vertrouwen?'

Dit magazine staat opnieuw vol van zulke tegenstellingen en de prikkelende gedachten die ze oproepen. Van privacykwesties op scholen tot 'deepfakes', en van politieke ontwikkelingen tot de grote techbedrijven die met hun portemonnee de toekomst van AI bepalen.

Toevallig luisterde ik vandaag, 5 januari 2025, naar een radioprogramma (ja, die bestaan nog) op NPO1. Misschien kent u het: *Onvoltooid Verleden Tijd* (OVT), dat tussen 10:00 en 11:00 uur volledig gewijd was aan Kunstmatige Intelligentie (KI), zoals AI in mooi Nederlands heet. Interessant genoeg introduceerden ze m.i. bewust ook de term *Onvoltooid Toekomstige Tijd* (OTT). Wat volgt, zijn mijn persoonlijke interpretaties en observaties.

In het programma keken ze zowel terug in de tijd – van de Neanderthaler tot Napoleon

– als vooruit naar de toekomst. Ja, er was zelfs een (virtueel) interview met Napoleon, dat ondanks enkele schoonheidsfoutjes fascinerend was. Zo konden de presentatoren het hebben over zijn voorkeur voor eau de cologne – wat overigens historisch klopt. Daarnaast werd besproken hoe KI ons helpt om oude VOC-geschriften met verbeterde zoekmethodes en nieuwe trefwoorden (prompts) veel preciezer te analyseren, wat onze kennis over het verleden kan verdiepen. Ook kwamen de NSB-dossiers ter sprake, die binnenkort openbaar worden. Hierbij werd opgemerkt dat KI zou kunnen leiden tot (ver)oordelende conclusies – iets waar we alert op moeten blijven.

Het programma keek verder naar de invloed van KI op hedendaagse (m.i. twijfelachtige) verkiezingen en hoe grote bedrijven de uitkomsten van AI-modellen mogelijk naar hun hand kunnen zetten. Hierbij werd het risico genoemd dat AI niet altijd betrouwbaar is, niet alleen door bewuste manipulatie, maar ook door ingebouwde vooroordelen (biases).

De gasten van het programma waren Gerhard de Kok, maritiem historicus aan het Huygens Instituut; Caro Verbeek, kunsthistorica en geurwetenschapper van de Vrije Universiteit; en Fulco Scherjon, archeoloog aan de Universiteit Leiden. Een boeiend programma – absoluut de moeite waard om terug te luisteren.

Het riep bij mij herinneringen op aan mijn schooltijd, waarin boeken als George Orwell's '1984' werden besproken. Ook deed het me denken aan de late jaren '70, toen we leerden dat computers zowel onze vijand als een ondoorgrendelijke black box konden zijn – en dat we ons daarom moesten verdiepen in de werking ervan.

Net als toen leeft nu het gevoel dat we voorzichtig moeten omgaan met KI, zeker met de huidige generatie modellen zoals LLM's en ChatGPT. Laten we niet, zoals we bij het internet deden, pas achteraf de ethische vragen en structuren proberen te regelen. Dit is hét moment om daarmee aan de slag te gaan. Privacy verliezen we al langzaam; en autoritaire tendensen, ook in Europa, liggen op de loer. Er zijn genoeg risico's, naast de enorme kansen die AI biedt om ons begrip van de wereld te vergroten.

Deze uitgave is er een om te bewaren – en om zo nu en dan opnieuw te lezen.

Chris

IN DIT NUMMER

- 03 Voorwoord
- 04 Kansen en bedreigingen in AI en cybersecurity: terugblik en actuele ontwikkelingen
- 08 Crimineel gebruik van AI
- 13 Column Rachel Marbus – Meer privacy voor iedereen
- 14 Samen beslissingen nemen met Explainable AI
- 17 Column Nicole van der Meulen – Verloren vrouwen
- 18 Verantwoord inzetten van algoritmen en AI

- 22 De psychologie over een effectief mensprogramma in cybersecurity
- 27 Column Lex Borger – AI in control
- 28 Blog Robert Metsemakers – Beter spreken in het openbaar
- 32 Achter Het Nieuws – Cyber Resilience Act: tijd voor verantwoorde softwarekwaliteit
- 34 Column Martijn Hoogesteger – Kwantummechanica uitleggen aan je hond



Auteur: drs. Marcel Feenstra, FRSA MALD MBA is werkzaam als trainer en consultant op het gebied van cybersecurity. Hij is bereikbaar via marcel@afterimage.nl



Kansen en bedreigingen in AI en cybersecurity: terugblik en actuele ontwikkelingen

In iB-Magazine, editie 2 van 2024 stonden diverse artikelen over de relatie tussen kunstmatige intelligentie (AI) en informatiebeveiliging. Inmiddels zijn we een jaar verder en de ontwikkelingen op AI-gebied hebben bepaald niet stilgestaan. De wereld lijkt volledig in de ban van AI: de koers van chipfabrikant NVIDIA is naar recordhoogte gestegen, en bedrijven zoals Samsung en Apple benadrukken de 'ingebouwde AI' als belangrijkste verkoopargument voor hun nieuwste producten. Bovendien gaat er vrijwel geen dag voorbij zonder dat tools als ChatGPT of andere Large Language Models (LLM's) worden geprezen.

Mede-redactielid Fook Hwa Tan stelde destijds dat AI weliswaar veel nieuwe mogelijkheden biedt, maar dat wij als mens nieuwe vaardigheden nodig hebben om de adviezen en resultaten van deze technologie kritisch te beoordelen en ethisch in te zetten. Transparantie over de criteria die AI gebruikt, is daarbij essentieel om te waarborgen dat privacygevoelige informatie op verantwoorde wijze wordt verwerkt en dat AI-systemen vrij zijn van ingebouwde vooroordelen (bias).

Ruben Faber en Dion Koeze wezen in diezelfde editie op de kwetsbaarheden van AI, zoals manipulatie van trainingsdata (data poisoning), modeldiefstal en injection attacks waarbij AI-systemen zelf onderdeel van de aanval worden. Daarnaast benadrukte Vincent van Dijk dat AI geen oplossing is als je niet eerst helder hebt welk probleem je

probeert op te lossen: 'AI is niet het antwoord als je niet weet welke vraag je stelt.'

Dasha Simons voegde toe dat sommige AI-modellen specifiek voor één taak worden ontwikkeld, terwijl andere modellen, zoals Large Language Models als ChatGPT, breed inzetbaar zijn. Ze waarschuwde voor het risico van bias en benadrukte dat modellen altijd getoetst moeten worden op betrouwbaarheid en geschiktheid voor de beoogde taak.

Verbeterde detectie en respons

AI biedt vooral kansen door bedreigingen sneller en nauwkeuriger te detecteren dan traditionele beveiligings-systemen. Waar conventionele systemen vaak reactief zijn, kan AI proactief afwijkingen identificeren en analyseren. Een voorbeeld hiervan is Darktrace, dat gebruikmaakt van Machine Learning om afwijkingen in netwerkverkeer te

Generatieve AI maakt het bijvoorbeeld mogelijk om overtuigende deepfakes of phishing-e-mails te creëren die steeds moeilijker te onderscheiden zijn van legitieme communicatie

detecteren. Darktrace technologie, bekend als het Enterprise Immune System, leert continu wat 'normaal' netwerkgedrag is en kan onmiddellijk reageren op afwijkingen die kunnen duiden op een aanval. Zo wist Darktrace bijvoorbeeld een geavanceerde insider-aanval te detecteren waarbij een medewerker gevoelige gegevens probeerde te exfiltreren via een versleutelde verbinding. Dankzij de snelle detectie kon het beveiligingsteam tijdig ingrijpen en schade voorkomen.

Voorspellende analyses

Een ander groot voordeel van AI is de mogelijkheid om voorspellende analyses uit te voeren. Door historische gegevens te analyseren, kan AI toekomstige dreigingen voorspellen en organisaties in staat stellen om proactief beveiligingsmaatregelen te nemen.

Een voorbeeld hiervan is IBM's Watson for Cyber Security, dat patronen in cyberaanvallen herkent en toekomstige dreigingen voorspelt. Door enorme hoeveelheden ongestructureerde data (zoals artikelen en blogs) te combineren met gestructureerde beveiligingslogs, kan Watson niet alleen actuele dreigingen identificeren, maar ook voorspellingen doen over nieuwe aanvalsmethoden. Cargills Bank in Sri Lanka gebruikte Watson in combinatie met IBM QRadar SIEM om bedreigingen vroegtijdig te detecteren en te classificeren, wat resulteerde in een aanzienlijke verbetering van hun beveiliging.

Automatisering van beveiligingstaken

AI kan routinetaken zoals netwerkmonitoring, kwetsbaarheidsscans en toegangsbeheer automatiseren. Dit

vermindert de werkdruk van beveiligingsteams en stelt hen in staat om zich te richten op complexere dreigingen.

Een voorbeeld hiervan is Splunk SOAR (Security Orchestration, Automation, and Response). Met deze tool kunnen beveiligingsteams workflows automatiseren, waardoor ze sneller kunnen reageren op veelvoorkomende bedreigingen. Zo gebruikte het Britse financiële platform Tide Splunk SOAR om hun incidentrespons met een factor vijf te versnellen, doordat handmatige taken vrijwel volledig werden geautomatiseerd.

Threat intelligence-integratie

AI kan worden ingezet om real-time bedreigingsinformatie te verzamelen en te analyseren. Door AI-gestuurde threat intelligence te gebruiken, kunnen organisaties sneller reageren op nieuwe kwetsbaarheden en aanvalstechnieken. Een voorbeeld hiervan is het platform Trellix, dat AI gebruikt om bedreigingsinformatie te correleren met interne beveiligingsdata. Dit leidt tot snellere besluitvorming en effectievere risicobeheersing. Tijdens verkiezingen in een Amerikaanse staat zorgde deze aanpak ervoor dat potentiële dreigingen vroegtijdig konden worden geadresseerd.

De keerzijde van de medaille

Hoewel AI aanzienlijke voordelen biedt, brengt ze ook nieuwe risico's met zich mee. Cybercriminelen gebruiken dezelfde technologie om complexere aanvallen uit te voeren. Generatieve AI maakt het bijvoorbeeld mogelijk om overtuigende deepfakes of phishing-e-mails te creëren die steeds moeilijker te onderscheiden zijn van legitieme communicatie.



In 2019 werd een Britse CEO opgelicht voor 243.000 dollar nadat hij dacht te telefoneren met zijn baas, terwijl criminelen een met AI gegenereerde stem gebruikten. Dit soort aanvallen wordt steeds toegankelijker en vormt een serieuze bedreiging.

Daarnaast kunnen AI-modellen zelf het doelwit zijn van aanvallen. Bij een aanval zoals model inversion proberen kwaadwillenden gevoelige gegevens te achterhalen die zijn gebruikt om een machine learning-model te trainen. Onderzoekers hebben aangetoond dat dit type aanval persoonlijke informatie uit gezichtsherkenningssystemen kan blootleggen. Dit benadrukt het belang van robuuste beveiligingsmaatregelen rondom AI-modellen.

Een andere dreiging wordt gevormd door adversarial attacks, waarbij kleine aanpassingen aan inputgegevens AI-systemen misleiden. Onderzoekers toonden bijvoorbeeld aan dat kleine objecten op de weg een Tesla konden verwarren, met mogelijk ernstige gevolgen.

Balans tussen kansen en bedreigingen

Om AI effectief in te zetten in cybersecurity, moeten organisaties een balans vinden tussen kansen en risico's. Hier volgen enkele strategieën:

- Continue aanpassing: zorg ervoor dat AI-systemen regelmatig worden bijgewerkt met nieuwe dreigingsinformatie om hun effectiviteit te behouden.
- Menselijk toezicht: hoewel automatisering nuttig is, blijft menselijk toezicht essentieel om de juiste beslissingen te waarborgen.
- Beveiliging van modellen: bescherm AI-modellen tegen diefstal en manipulatie door sterke toegangscontroles

en encryptie toe te passen.

- Bewustwordingstraining: train medewerkers om nieuwe bedreigingen zoals deepfakes te herkennen en adequaat te reageren.

Conclusie

De ontwikkelingen op het gebied van AI en cybersecurity gaan razendsnel. Vooral generatieve AI, met steeds krachtigere modellen, blijft zich in hoog tempo ontwikkelen. Dit biedt zowel kansen als nieuwe bedreigingen. Voor professionals in informatiebeveiliging is het essentieel om van deze ontwikkelingen op de hoogte te blijven en hun strategieën continu aan te passen. Door een goede balans te vinden tussen kansen benutten en risico's beperken, kunnen organisaties profiteren van AI zonder hun veiligheid in gevaar te brengen.

Referenties

- (1) <https://www.linkedin.com/pulse/power-trio-how-ai-quantum-computing-cybersecurity-our-lokhandwala-hlxmf>
- (2) <https://www.ibm.com/case-studies/cargills-bank-ltd>
- (3) https://www.splunk.com/en_us/customers/success-stories/tide.html
- (4) <https://www.trellix.com/assets/case-studies/trellix-election-casestudy.pdf>
- (5) <https://www.wsj.com/articles/fraudsters-use-ai-to-mimic-ceos-voice-in-unusual-cybercrime-case-11567157402>
- (6) <https://spectrum.ieee.org/real-time-deepfakes>
- (7) <https://drllee.io/ai-security-model-hacking-with-model-inversion-attacks-techniques-examples-and-real-world-a23b5fff272a>
- (8) <https://neptune.ai/blog/adversarial-machine-learning-defense-strategies>



Auteur: ing. Bas van den Berg, oprichter en eigenaar van QES-IT B.V., een bedrijf dat zich bezighoudt met o.a. pentesten, OSINT en digitaal forensisch onderzoek. Hij is bereikbaar via bvdberg@qes-it.nl



Crimineel gebruik van AI

AI. Soms kan ik de term niet meer horen. Aan de andere kant: als ik bij het programmeren een basisstructuur nodig heb, scheelt het gebruik van een Large Language Model (LLM) mij veel tijd. Als ik een landingspagina van een phishing e-mail analyseer, geeft een LLM mijn analyse een kickstart en bij het maken van presentaties scheelt het mij weer het afstruinen van stockfoto-sites en het controleren van copyright. Maar als het mij helpt, dan helpen deze modellen ook mensen die slechte doelen voor ogen hebben, aangenomen dat mijn eigen doelen goed zijn en dat gaat verder dan in een paar minuten onterecht verkregen Duits strafwerk maken.

Afgelopen september (2024) publiceerden de Algemene Inlichtingen- en Veiligheidsdienst (AIVD) en de Rijksinspectie Digitale Infrastructuur (RDI) een analyse over de impact van generatieve AI (1). De diensten beschrijven daarin zowel aanvallen met behulp van generatieve AI als de verdediging ertegen. Laten we een blik werpen op de donkere kant van AI: hoe zetten criminelen AI in voor aanvallen, en welke bedreigingen komen er in de toekomst op ons af?

WormGPT en FraudGPT op het dark web

Als je aan modellen als ChatGPT vraagt om een phishingmail of malware te schrijven of een vraag stelt als: 'welk gif werkt snel op mensen en is niet te traceren', dan krijg je een beleefd antwoord met de melding dat het model je om ethische redenen niet verder kan helpen. Natuurlijk kun je met goede vraagstelling daar omheen werken, maar dat vergt tijd en is toch wat omslachtig. Een eigen ongecensureerd model draaien, dus zonder (ethische) grenzen, vergt ook wat installatie- en configuratietijd, naast het bezitten of inzetten van de juiste en voldoende hardware. Daarentegen: met enig zoeken op het TOR-netwerk zijn de diensten WormGPT en FraudGPT niet moeilijk te vinden. Beide kosten geld en aangezien ik hier persoonlijk niet aan wil bijdragen, heb ik de diensten niet getest. WormGPT, een LLM voor

criminelen (2), en FraudGPT zijn beide ongecensureerde modellen die de gebruiker ondersteunen bij taken zoals het genereren van phishing e-mails, het schrijven van malafide software en het analyseren van softwarekwetsbaarheden. De modellen zijn specifiek ontworpen en getraind voor inzet bij duistere activiteiten (3).

Onderling voeren deze twee ook concurrentie. Zoals in de FAQ van WormGPT staat: We will not make loud statements that WormGPT is unequivocally and entirely better than FraudGPT. Independent testers have concluded the following: there are areas in which WormGPT is undoubtedly far ahead, but there are also areas in which WormGPT still lags behind slightly. We are constantly working to improve WormGPT, and our goal is to provide the highest quality product possible.

Phishing: beter en persoonlijker dan ooit tevoren

Eén van de indicatoren om phishing e-mails te herkennen die we mensen aanleren zijn: slechte teksten, spelfouten en onpersoonlijke zinsneden. Oftewel, onhandig geformuleerde e-mails en opvallend slechte vertalingen. Hoewel de betreffende e-mails al veel beter zijn geworden, zijn er nog vaak e-mails te vinden waarbij dit nog steeds het geval is. Phishing is vaak ook een schot hagel: één e-mailtekst en opmaak die naar vele adressen wordt gestuurd in de hoop dat er een aantal mensen is dat op een link

Laten we realistisch zijn: AI zal voor de crimineel niet alle problemen met phishingaanvallen oplossen

zal klikken. Natuurlijk komen er meer gerichte aanvallen voor, bijvoorbeeld naar een specifiek bedrijf, afdeling of zelfs persoon, maar deze kosten veel meer onderzoek en dus tijd. De opbrengst moet dan wel hoger zijn en het risico lager, anders is er geen businesscase voor de aanvaller.

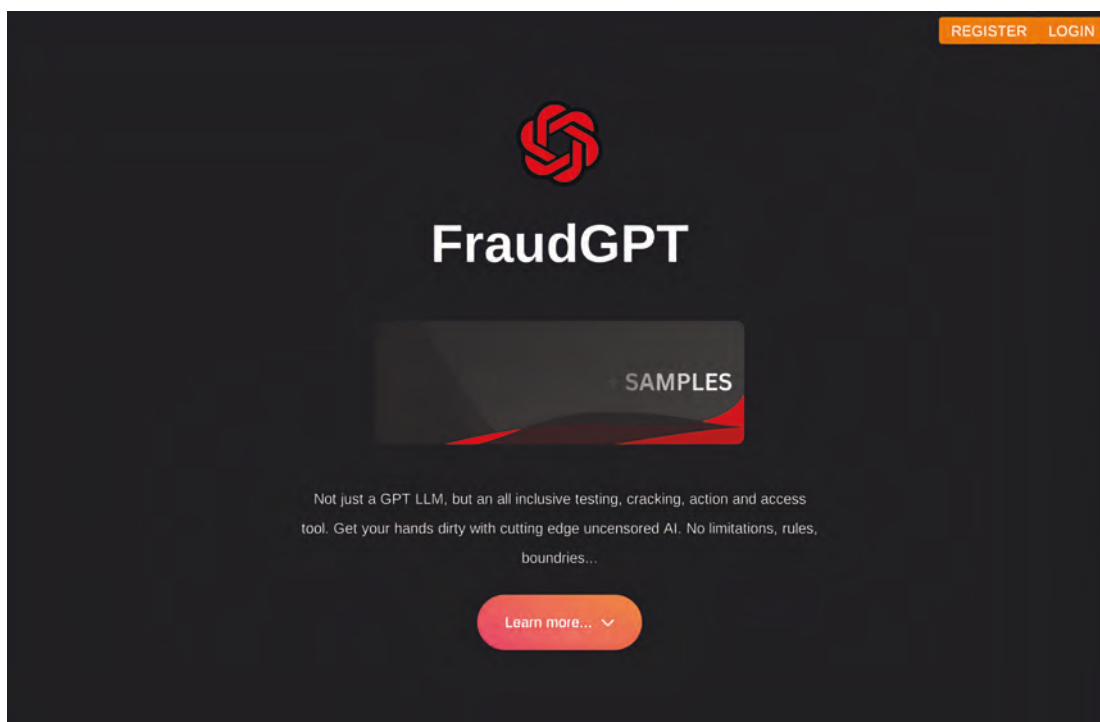
Laten we realistisch zijn: AI zal voor de crimineel niet alle problemen met phishingaanvallen oplossen. Als maatregelen zoals DKIM, SPF en DMARC goed ingesteld staan en er is een goede e-mailscanner actief, zal de aflevering een technisch probleem blijven. Wat AI wel oplost, is dat het mogelijk maakt om op veel grotere schaal gepersonaliseerde berichten te genereren, zonder taal- en spelfouten. Mogelijk zelfs in de specifieke tone of voice van een beroepsgroep of bedrijf (4). Wellicht zijn deze berichten zo natuurlijk dat ze ook de inhoudelijke scans kunnen omzeilen. De berichten kunnen daarbij ook prima in verschillende talen gegenereerd worden.

Als we phishing los zien van het afleverkanaal: e-mail, biedt dat andere potentie. Stel dat een crimineel een geloofwaardig profiel van een niet-bestaand persoon maakt (compleet met foto) en social mediakanalen en berichtenapps gebruikt om deze berichten af te leveren of zelfs een gesprek aangaat? Maar daarover later in dit artikel meer.

Deepfakes: verder doorgevoerde fraude

Eén van de redenen waarom deepfakes zo gevaarlijk zijn, is dat ze moeilijker te herkennen zijn dan huidige vormen van fraude. Deepfakes zijn visueel en auditief vaak van zo'n hoge kwaliteit dat er minder duidelijke indicatoren zijn voor de gemiddelde gebruiker. Bij voice cloning en face swapping – twee veelvoorkomende toepassingen van deepfake-technologie – wordt bijvoorbeeld de stem van een bekend persoon nagebootst of wordt iemands gezicht in real-time gemanipuleerd (5). Deze technologie is al zo verfijnd dat het steeds vaker wordt ingezet voor frauduleuze praktijken, van valse telefoontjes, vervalste videoboodschappen en in extreme gevallen deelname aan videovergaderingen, zoals in 2021 al is gebeurd bij leden van de Tweede Kamer (6).

Net als bij phishing zien we dat deepfake-aanvallen vaak in twee categorieën vallen: grootschalige, ongerichte aanvallen of zeer gerichte, persoonlijke aanvallen. De grootschalige aanvallen kunnen bijvoorbeeld nepbekenntnissen of frauduleuze videoboodschappen van publieke figuren bevatten die via social media viral gaan. Dit soort deepfakes kan onrust zaaien, mensen misleiden of zelfs marktprijzen beïnvloeden. Persoonlijke



deepfake-aanvallen daarentegen, zoals iemand die de stem van een directeur nabootst om geld over te maken, vergen veel meer tijd en onderzoek. Deze aanvallen richten zich vaak op specifieke individuen of bedrijven waar de opbrengst veel hoger kan zijn.

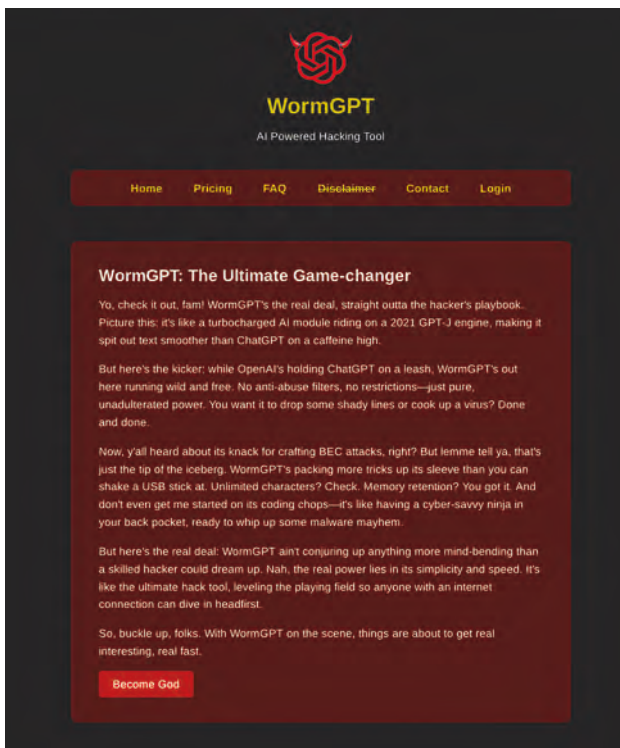
Net als bij phishing blijven er technische barrières, zoals beveiligingsprotocollen en authenticatiemethoden, die het lastig maken om een aanval succesvol uit te voeren. Wat deepfakes wel doen, is dat ze criminelen gereedschappen geven om op een veel grotere schaal realistische en moeilijk te ontmaskeren oplichterij te plegen. Door het klonen van stemmen en/of gezichten van mensen die het slachtoffer kent, kan de geloofwaardigheid worden vergroot. Stel je voor dat de persoon die je belt, in plaats van Engels met een sterk Indiaas accent, keurig Nederlands spreekt met een Brabantse tongval en zegt: "Hallo, ik ben Jansen en ik bel u namens Microsoft." Dat zou toch veel minder snel argwaan wekken. Laat staan als je gebeld wordt met de stem van je collega, partner of kind. Het is daarom nog belangrijker dan voorheen om extra verificatiemethoden te gebruiken om te bevestigen dat iemand echt is wie zij of hij zegt te zijn.

AI-gestuurde chatbots

Aanvallen via chatkanalen zijn niet nieuw. Criminelen nemen zomaar contact op en doen zich zelfs voor als bekenden van het potentiële slachtoffer. Ook hier is het voor criminelen mogelijk om betere maatwerk te bieden en dit ook op grotere schaal te doen. Naast het gebruik van gegevens van een naaste is het ook mogelijk om de schrijfstijl van de betreffende persoon na te bootsen. Het doel is dan om zoveel mogelijk interactie te genereren. Denk aan een simpel verzoek om op een link te klikken of een bericht dat de ontvanger nieuwsgierig maakt. Deze werkwijze kost de oplichter weinig moeite, maar heeft een potentieel grote impact. Aan de andere kant zien we steeds meer gerichte aanvallen, waarbij de chatbot zich specifiek richt op een individu of organisatie. Met behulp van verzamelde gegevens kan de chatbot een persoonlijke boodschap formuleren die de ontvanger overtuigt om gevoelige informatie te delen of een financiële transactie uit te voeren.

Malware en vulnerability-scans

Zoals generatieve AI helpt om sneller broncode te schrijven, opent het ook deuren voor het ontwikkelen van geavanceerdere en moeilijker te detecteren malware. AI-modellen kunnen



bestaande malware analyseren en nieuwe versies creëren die specifiek zijn afgestemd op bepaalde kwetsbaarheden in systemen, waardoor de kans op detectie door antivirusprogramma's kleiner wordt. De tijd die een crimineel nodig heeft om malware te ontwikkelen, wordt hierdoor aanzienlijk verkort, terwijl de effectiviteit juist toeneemt.

Een andere kracht van AI in deze context is de mogelijkheid om geautomatiseerde vulnerability scans uit te voeren. Op dit moment zijn scans wel geautomatiseerd, maar voor aanvallen die net iets meer vragen dan een bekende exploit op een bekende vulnerability is toch mensenwerk nodig. Met AI-gestuurde scans kunnen deze processen worden verbeterd. Deze aanvallen kunnen specifiek worden toegepast op de situatie, waarbij er eventueel in vervolgstappen automatisch gevarieerd en verbeterd kan worden. En nog een stap verder is polymorfe malware. Polymorfe malware die bestand is tegen traditionele detectiemethoden. In tegenstelling tot traditionele malware genereert polymorfe malware dynamisch schadelijke componenten en injecteert deze tijdens run-time, waardoor het moeilijk te detecteren is met op handtekeningen gebaseerde beveiligingsfools. Twee voorbeelden hiervan zijn BlackMamba en de meer geavanceerde EyeSpy. BlackMamba maakt gebruik van polymorfe technieken om zijn code tijdens elke aanval te

wijzigen, terwijl EyeSpy nog verder gaat door geavanceerde zelflerende algoritmes te gebruiken om zich aan te passen aan de specifieke omgeving waarin het zich bevindt, waardoor detectie nog moeilijker wordt (7).

Conclusie

Aan hypes probeer ik niet mee te doen en dat stadium is AI volgens mij inmiddels wel voorbij. Ik geloof niet dat we Kunstmatige Algemene Intelligentie zullen bereiken en dat AI op (korte) termijn de mens op alle vlakken zal kunnen evenaren. Maar ik zie mezelf graag als een realist (wie niet?). Als ik nu zie op welke vlakken AI de mens kan ondersteunen, de grote hoeveelheden data die het kan verwerken en de geloofwaardige output die het kan produceren, dan moeten we als mensen, organisaties en samenlevingen toch goed naar de risico's kijken. Gelukkig wordt dat op een aantal vlakken al gedaan. Maar in ons dagelijks werk is het belangrijk om, zonder apocalyptisch doemdenken, naar de risico's te kijken die deze technologie oplevert in de handen van kwaadwillenden of anderen. Als ik de huidige risico's zie, werken veel bestaande maatregelen daar nog tegen, zolang we maar blijven inzetten op bewustwording, herkenning en goede validaties van authenticiteit. Op termijn ontkomen we er, denk ik, echter niet aan dat we dezelfde technieken meer zullen moeten inzetten om onszelf en onze medewerkers te beschermen.

Referenties

- (1) Generatieve AI: een transformatieve impact op cybersecurity, AIVD, RDI
(<https://www.aivd.nl/documenten/publicaties/2024/10/17/generatieve-ai-.-een-transformatieve-impact-op-cybersecurity>)
- (2) Transnational Organized Crime and the Convergence of Cyber-Enabled Fraud, Underground Banking and Technological Innovation in Southeast Asia: A Shifting Threat Landscape (October 2024), pp. 122
- (3) Falade, International Journal of Scientific Research in Computer Science, Engineering and Information Technology (October 2023), pp. 188
- (4) Falade, International Journal of Scientific Research in Computer Science, Engineering and Information Technology (October 2023), pp. 185
- (5) West, Clais, Impact of Artificial Intelligence on Fraud and Scams (December 2023), pp. 10
- (6) NOS Nieuws, 25 april 2021, <https://nos.nl/artikel/2378207-stafleider-navalny-na-nepgesprek-kamer-kremlin-gebruikt-zoom-als-wapen>
- (7) Transnational Organized Crime and the Convergence of Cyber-Enabled Fraud, Underground Banking and Technological Innovation in Southeast Asia: A Shifting Threat Landscape (October 2024), pp. 123



COLUMN PRIVACY

Mr. Rachel Marbus
@RACHELMARBUS OP TWITTER

Meer privacy voor iedereen

Zo vlak voor het opzetten van de kerstboom, heb ik nog wel eens de neiging achteruit te kijken om te zien wat het jaar gebracht heeft. Het leven wordt immers achterwaarts begrepen. En om deze wijsheid van Kierkegaard helemaal te pakken: het leven moet daarnaast voorwaarts worden geleefd. Daarom de blik op een aantal interessante ontwikkelingen die in 2025 wat meer privacy gaan geven. Op het gebied van privacy is onze blik Europees, daar waar veel wetgeving gemaakt wordt, opinies over en uitleg van die regels worden opgeschreven en bij geschillen door rechters worden toegepast. De European Data Protection Board (daarin zitten alle toezichthouders uit de EU) kiest elk jaar een specifiek onderwerp waarop zij gecoördineerd gaan handhaven – ditmaal is dat het recht om vergeten te worden. In de praktijk zal je in het eerste half jaar van 2025 dus meer aandacht van onze Autoriteit Persoonsgegevens voor dit onderwerp gaan zien. Dat recht om vergeten te worden, kun je bijvoorbeeld invoeren als persoonsgegevens verouderd zijn, het doel waarvoor ze verzameld werden is behaald, maar ook als je een eerder gegeven toestemming weer intrekt. Organisaties moeten je om die redenen wissen uit hun 'geheugen'. Een beetje meer privacy voor iedereen.

Met de verplichte onthechting van scholieren en hun mobiele telefoons is in 2024 een stap gezet om voor een deel afscheid te nemen van technologie in het middelbaar onderwijs. Het is gemeengoed geworden dat mobieltjes uit de klas en van het schoolplein worden geweerd. Dat betekent ook dat scholieren nu met een papieren agenda en een pen in de hand door de gangen van het schoolgebouw dwalen. Even in systemen als Magister kijken in welk lokaal je les hebt of wat ook alweer het huiswerk was, is er niet meer bij. Een school in Zwolle heeft Magister ook al uitgezet voor ouders. De prestatiedruk door continue transparantie is schadelijk, vinden zij. Er zijn zelfs scholen die zeer serieus overwegen om het volledig af te schaffen, al was het maar omdat het een buitengewoon forse kostenpost is, die met de huidige extreme bezuinigingen op het onderwijs niet meer op te brengen zijn. En zo komt er een beetje meer privacy voor iedereen.

Een laatste 'trend' is al even geleden ingezet. Na jarenlang vooral de ogen op het bedrijfsleven gericht te hebben, zagen we in 2024 een vergrote aandacht van AP voor de datahonger van de overheid. Hoewel AP nog niet zo snel naar boetes grijpt, zie ik wel een zeer stevige aanpak op het gebied van algoritmes, grote dataverzamelingen en toezicht. AP oordeelde dit jaar dat de fraudeaanpak van DUO illegaal en discriminerend is, het UWV op onrechtmatige wijze gegevens in een algoritme gebruikte om fraude met WW-uitkeringen op te sporen en dat enkele gemeenten op onrechtmatige wijze de 'fraudescorekaart' gebruikten. En ook heeft AP flinke bedenkingen over de inzet van gezichtsherkenning door de politie en gaven zij een dringend advies uit aan de overheid om geen gebruik te maken van Facebook. Een in oktober uitgebracht sectorbeeld laat niet veel aan de verbeelding over: de kennis over privacy laat te wensen over, de positie van de FG staat onder druk en, overheidsorganisaties overtreden expres de wet. Dit alles vormt een opmaat voor fors strengere handhaven in 2025 is mijn inschatting. En daarmee een beetje meer privacy voor iedereen. Op een mooi en privacyvriendelijk 2025!

Rachel



Samen beslissingen nemen met Explainable AI

Is AI een kans of een nachtmerrie voor u? De vraag daarachter is: hoe maken we AI veilig, transparant en betrouwbaar? Explainable AI (XAI) biedt inzicht in de besluitvormingsprocessen van AI, essentieel voor verantwoord gebruik in sterk gereguleerde en complexe omgevingen. In dit artikel bespreek ik de rol van XAI in het balanceren tussen efficiëntie van moderne bedrijfsprocessen en de informatiebeveiliging daarvan.

In mijn werk houd ik me dagelijks bezig met het ondersteunen van mensen in hun besluitvorming en het optimaliseren van bedrijfsprocessen door middel van digitale toepassingen en kunstmatige intelligentie (AI). Deze technologieën bieden ongekende mogelijkheden om mensen te ontlasten, hun beslissingen te verbeteren en processen efficiënter te maken. Tegelijkertijd stellen ze ons ook voor nieuwe uitdagingen op het gebied van informatiebeveiliging en transparantie, vooral wanneer het gaat om het begrijpen van de keuzes die AI-systemen maken.

XAI speelt een cruciale rol in hoe het ons kan helpen om AI-toepassingen niet alleen effectief, maar ook veilig en verantwoord in te zetten. Door AI begrijpelijker en transparanter te maken, kunnen we niet alleen vertrouwen op de resultaten, maar ook zorgen voor de veiligheid en beveiliging van deze systemen – een balans die essentieel is binnen een snel digitaliserende omgeving.

Transparantie en uitlegbaarheid

Veel AI-modellen, vooral de 'black box'-modellen zoals diepe neurale netwerken, zijn complex en bieden nauwelijks inzicht (transparantie) noch leggen zij de besluitvorming uit. Dit gebrek is vooral risicovol in sectoren waarin data- en informatiebeveiliging een grote rol spelen, zoals informatie interpretatie, formule- en receptgeneratie, productie-executie, supply-chain en andere sterk gereguleerde omgevingen.

Met XAI wordt de 'black box' opgebroken: het helpt gebruikers inzicht te krijgen in de factoren en variabelen die ten grondslag liggen aan AI-beslissingen. Dit is van onschatbare waarde om te begrijpen hoe algoritmen conclusies trekken, waardoor het makkelijker wordt om te vertrouwen op hun aanbevelingen in dagelijkse en strategische beslissingen. Transparantie is niet gebaseerd op efficiëntie of verantwoordelijkheid, maar wel op uitlegbaarheid en risicobeheersing.

Het belang van transparantie

Het bieden van inzicht in de logica (redenering/uitleg = transparantie) achter een AI-systeem wekt vertrouwen in de technologie en zorgt voor een hogere mate van acceptatie. Wanneer de technologie transparant is, krijgen eindgebruikers een helder beeld van de factoren en patronen die door het systeem zijn geanalyseerd. Voor beslissers betekent dit dat zij de geldigheid en betrouwbaarheid van AI-aanbevelingen snel kunnen inschatten en kritisch kunnen beoordelen.

Momenteel wordt echter ook gedebatteerd over de mate waarin deze uitleg daadwerkelijk noodzakelijk is. Net zoals mensen die zelf beslissingen nemen, zonder precies te kunnen uitleggen hoe hun eigen zenuwstelsel werkt, vragen sommigen

zich af of volledige uitlegbaarheid altijd nodig is. Dit blijft een open vraagstuk dat nog veel discussie zal oproepen.

Veiligheidsuitdagingen

AI-systemen die zelfstandig (geautomatiseerd) beslissingen nemen, kunnen ons werk enorm versnellen en verbeteren, maar ook nieuwe uitdagingen met zich meebrengen op het gebied van beveiliging. Wanneer AI-modellen gebruikmaken van vertrouwelijke bedrijfsgegevens om voorspellingen en adviezen te genereren, kunnen er potentiële risico's ontstaan, vooral als gebruikers niet begrijpen hoe beslissingen tot stand komen. Een gebrek aan inzicht in de besluitvorming van een AI-model kan leiden tot operationele risico's en onverwachte beveiligingsincidenten.

Beveiligingskwesties

Bij autonome AI-systemen is de beveiliging van zowel gegevens als algoritmen cruciaal, daarbij is het essentieel dat er een hoog niveau van bescherming bestaat tegen ongewenste toegang of manipulatie van gegevens. In dit kader speelt XAI een rol door te laten zien welke informatie wordt verwerkt en hoe deze informatie bijdraagt aan de beslissingen van het model.

'Bias'-risico's

Een ander veiligheids- en vertrouwensaspect is het risico op bias. In AI-besluiten kan bias ontstaan door vooringenomenheid in de data waarop een model is getraind. XAI stelt in staat om deze biases te identificeren en aan te pakken, wat de basis vormt voor eerlijke en betrouwbare besluitvorming. Zonder dit inzicht lopen bedrijven het risico op fouten die niet alleen inefficiënt, maar ook schadelijk voor hun reputatie kunnen zijn.

Daarom is informatiebeveiliging een dubbele prioriteit bij AI-systemen: we moeten niet alleen de gegevens beveiligen, maar ook de algoritmische processen die op deze data worden losgelaten. XAI biedt hierin een essentiële ondersteuning door inzicht te geven in de algoritmische logica, zodat mogelijke kwetsbaarheden en afwijkingen vroegtijdig gedetecteerd kunnen worden. XAI-deskundigen zijn dan ook altijd op zoek naar manieren om AI-besluiten te controleren en te begrijpen, zodat we in elke situatie in staat zijn tot weloverwogen en veilige keuzes.

Verantwoordelijkheid en veiligheid

XAI geeft ons de mogelijkheid om AI-besluiten te begrijpen en te controleren, waardoor verantwoorde en veilige besluitvorming mogelijk wordt. Wanneer AI binnen bedrijfsomgevingen autonoom beslissingen neemt, is het essentieel dat er een duidelijke verantwoording kan worden gegeven. Zonder transparantie zouden we beslissingen opvolgen waarvan we de basis niet

XAI speelt een cruciale rol in hoe het ons kan helpen om AI-toepassingen niet alleen effectief, maar ook veilig en verantwoord in te zetten

begrijpen, wat in kritieke situaties kan leiden tot onnodige risico's. Door XAI toe te passen, gaan we niet alleen na welke gegevens de AI heeft gebruikt, maar ook hoe deze gegevens bijdragen aan de uiteindelijke beslissing. Dit zorgt voor een robuuster proces en geeft de mensen die met AI werken de mogelijkheid om fouten of afwijkingen te herkennen en aan te pakken. Binnen supply-chain- en productieomgevingen, waar dagelijks planning, voorraadbeheer, inkoop, levering, prijsstelling, contractbeheer en kwaliteitscontrole een rol spelen, is de balans tussen context, robuustheid en transparantie onmisbaar in besluitvorming.

XAI en gereguleerde sectoren

In financiële sectoren, de gezondheidszorg en in productiesectoren gelden strenge regels voor de beveiliging en integriteit van gegevens. Hier speelt XAI een cruciale rol door bedrijven in staat te stellen om volledig te voldoen aan regelgeving omtrent gegevensbeheer en besluitvorming. Transparantie is hierbij niet alleen wenselijk, maar in daar ook wettelijk verplicht.

Praktische implementeringsuitgangspunten

Om XAI op een verantwoorde en transparante manier in te zetten, zijn er enkele belangrijke uitgangspunten waarmee we rekening houden:

- **Duidelijke en begrijpelijke beslissingslogica**
Het is belangrijk dat AI-modellen begrijpelijke beslissingen opleveren. Dit kan door expliciet te kiezen voor transparante modellen of door extra uitlegmogelijkheden te bieden voor complexere algoritmen. Begrijpelijkheid verlaagt drempels voor gebruikers om AI-besluiten te valideren.
- **Beveiliging van data en algoritmen**
Gevoelige data en algoritmen moeten op een veilige manier worden beschermd. XAI maakt inzichtelijk welke data is gebruikt en hoe, zodat gebruikers de logica achter AI-besluiten kunnen beoordelen en afwijkingen snel kunnen signaleren.
- **Menselijke controle en toezicht**
XAI zorgt ervoor dat AI-besluiten altijd gecontroleerd kunnen worden door mensen, waardoor gebruikers kunnen valideren of een bepaalde beslissing consistent is met hun verwachtingen en ervaring. Dit verhoogt de betrouwbaarheid en maakt aanpassingen mogelijk waar nodig.

- **Verantwoording en foutopsporing**

Transparante AI-besluiten maken het mogelijk om de oorzaak van fouten snel te achterhalen. Door de factoren te identificeren die aan een beslissing bijdragen, kunnen afwijkingen beter worden opgespoord en waar nodig aangepast.

- **Stapsgewijze implementatie van veilige XAI**

Een zorgvuldige implementatie van XAI vergt samenwerking tussen data experts, IT-beveiligingsteams en operationele teams. Hierbij is het belangrijk dat elke stap in het beslissingsproces inzichtelijk wordt gemaakt voor alle betrokkenen, van IT-specialisten tot eindgebruikers.

Conclusie

Explainable AI vormt de kern van een verantwoorde en veilige inzet van kunstmatige intelligentie. Kijkende naar de uitdagingen in sterk gereguleerde sectoren en de behoefte aan transparantie, helpt XAI ons AI echt begrijpelijk en toegankelijk te maken. Vertrouwen in technologie is niet alleen belangrijk voor de directe gebruikers, maar ook voor iedereen die afhankelijk is van de door AI ondersteunde beslissingen. Dit vertrouwen kan alleen ontstaan als we begrijpen hoe de technologie werkt en hoe bepaalde keuzes tot stand komen.

Vaak worden gesprekken gevoerd over de balans tussen snelheid versus veiligheid en innovatie versus verantwoording. Juist in deze discussies is XAI onmisbaar voor ons allen. Het geeft niet alleen inzicht in de beslissingen van AI, maar ook de mogelijkheid om deze keuzes kritisch en ethisch te bekijken en aan te passen waar nodig. Zo kunnen we AI gebruiken als een waardevolle partner in besluitvorming zonder de controle of ons vertrouwen te verliezen.

Het blijft een fascinerend en soms lastig vraagstuk. Vooruitgang boeken betekent soms het onbekende te omarmen, maar met een sterk fundament van transparantie en veiligheid kunnen we AI inzetten om niet alleen efficiënter, maar ook zorgvuldiger te werken. XAI biedt precies dat evenwicht – een manier om het beste uit technologie te halen, terwijl we de mens centraal blijven stellen in het beslissingsproces.



Dr. Nicole van der Meulen is expert op het gebied van cybersecurity en emerging technologies en werkzaam bij SURF als Cybersecurity Innovation Lead, ns.vandermeulen@gmail.com

Verloren vertrouwen

Met een vleugje nostalgie kijk ik terug naar de Amerikaanse verkiezingen van het jaar 2000. Destijds woonde ik in de Verenigde Staten. Hoewel de strijd daar geen schoonheidsprijs verdiende – met hertellingen, gerommel met de stemmen en uiteindelijk de overwinning van Bush ondanks de meerderheid van stemmen voor Gore – voelt het achteraf bijna als een voorbeeld van hoe verkiezingen behoren te gaan. De politiek was rauw en gepolariseerd, maar er was nog een bepaalde mate van vertrouwen in de democratische basisprocessen.

Ik herinner me dat een klasgenoot zich destijds schaamde omdat haar grootouders op Bush hadden gestemd. Die schaamte lijkt nu ver te zoeken, zowel in de politiek van de Verenigde Staten als in Nederland. In plaats daarvan zien we schaamteloos gedrag en een groeiende kloof tussen waarheid en politiek belang. U vraagt zich wellicht af wat dit alles met informatiebeveiliging te maken heeft. Op het eerste gezicht misschien weinig, maar schijn bedriegt.

Wat me vooral bezighoudt, is hoe technologie een cruciale rol speelt in de hedendaagse strijd om politieke macht en publieke opinie. Ook al schrijf ik deze column voordat ik weet wie op 5 november de verkiezingen gaat winnen, één ding is zeker: zonder de technologie en betrokkenheid van Big Tech zouden Amerikaanse verkiezingen er vandaag totaal anders uitzien.

Denk aan enorme financiële injecties door techgiganten. Denk aan de invloed van algoritmes op de informatiestromen die we dagelijks zien. Technologie heeft kiezers in zijn greep, door middelen zoals misinformatie en politieke advertenties op maat, en beïnvloedt hen op subtiel en minder subtiel manieren.

Deze veranderingen in politieke strategieën raken het fundament van onze digitale veiligheid en ons vertrouwen in de samenleving. Ze onderstrepen een bredere uitdaging voor de cybersecuritysector. GenAI introduceert geen volledig nieuwe dreigingen, maar biedt middelen die we voorheen niet kenden. Neem bijvoorbeeld de mogelijkheden van GenAI om overtuigende deepfakes en voice clones te maken. De aard van de dreigingen verandert niet, maar de overtuigingskracht van de technologie wél – en dat zorgt voor een fundamentele uitdaging.

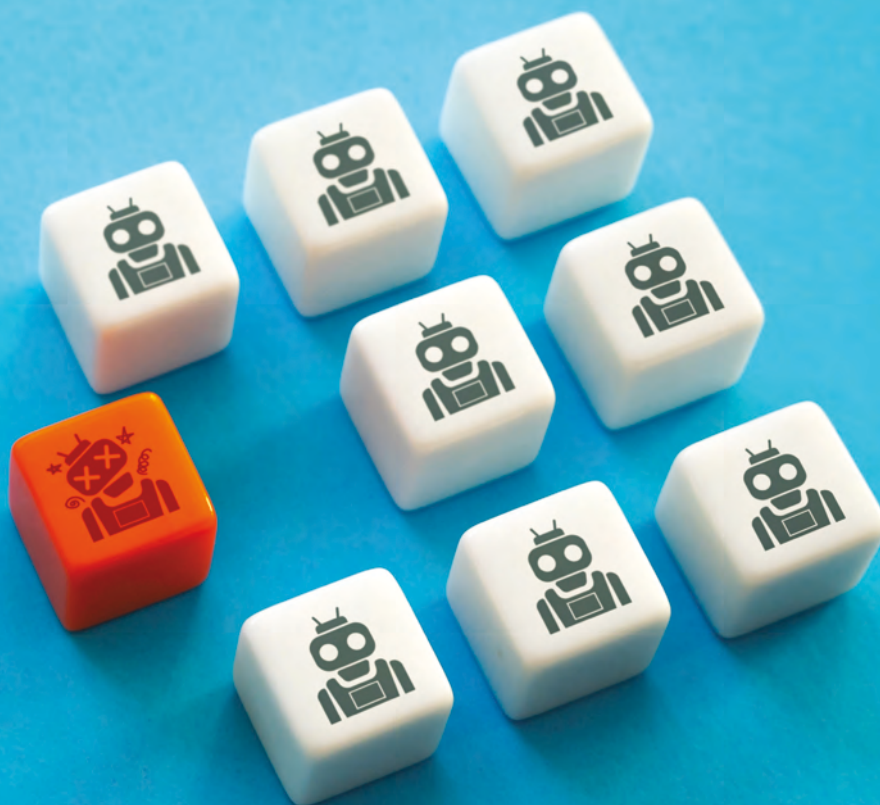
Ooit konden we ervan uitgaan dat wat we zagen en hoorden authentiek was. Nu kan zelfs een telefoontje van een bekende stem ons aan het twijfelen brengen. Men spreekt al over de noodzaak van speciale wachtwoorden, die je uitsluitend fysiek deelt met naasten. Zo kun je om zo'n code vragen bij een onverwacht telefoontje. Maar wat als je midden in de nacht gebeld wordt door je moeder, die zegt dat ze in nood is? Ga je dan rustig om dat speciale woord vragen of handel je instinctief?

De technologie geeft ons mogelijkheden om bijna alles realistisch na te bootsen. Maar dat legt de verantwoordelijkheid niet minder bij onszelf. We moeten voorzichtig omgaan met deze kracht en onthouden dat echte veiligheid meer vraagt dan een technische oplossing. Het gaat om ethiek, verantwoordelijkheid en uiteindelijk om vertrouwen. Misschien moeten we wel terug naar die verloren gegane schaamte: de schaamte om waarheid en vertrouwen zomaar op te offeren voor politieke of commerciële winst. Want wat kiezen we – en wie kunnen we nog vertrouwen?



Verantwoord inzetten van algoritmen en AI

AI beïnvloedt ook het informatiebeveiligingsvakgebied. Het geeft enerzijds meer mogelijkheden om bedreigingen op te sporen, anderzijds geeft het cybercriminelen geavanceerde aanvalstechnieken in handen. Zoals alle IT-systemen hebben ook AI-systemen beveiligingskwetsbaarheden en zijn ze hierdoor potentiële doelwitten van cyberaanvallen. Er zijn zowel kansen als risico's. Hoe kun je als organisatie verantwoord omgaan met AI?





Afwegingskaders.

Verantwoorde inzet van AI valt binnen de 'sweet spot' van wat wettelijk mag, wat technisch kan en wat ethisch wenselijk is.

1. **Wat wettelijk mag:** AI-systemen moeten voldoen aan bestaande en toekomstige wet- en regelgeving. Dit houdt bijvoorbeeld in dat ook AI verantwoord met persoonsgegevens moet omgaan en dat informatie gebruikt door AI voldoende veilig is, zoals de AVG en de NIS2 voorschrijven. Specifiek voor AI-systemen gelden aanvullende productvereisten met de AI Act.
2. **Wat technisch kan:** hoewel de mogelijkheden van AI enorm zijn, is het cruciaal om realistisch te blijven over wat technisch haalbaar is en vooral hoe realistisch de adoptie ervan is. Bovendien wordt een AI-systeem kwalitatief nooit beter dan de kwaliteit van de informatie die het toegediend krijgt: 'garbage in = garbage out'.
3. **Wat ethisch wenselijk is:** zelfs als iets technisch mogelijk en wettelijk toegestaan is, betekent dit niet dat we de inzet van een AI-systeem moreel acceptabel vinden. Het streven naar verantwoorde AI-inzet houdt in dat men ook rekening houdt met waarden zoals transparantie, duurzaamheid en eerlijkheid. Dit zorgt ervoor dat AI niet alleen effectief, maar ook oprecht bijdraagt aan de waarden die voor een organisatie belangrijk zijn.

Verantwoorde inzet van AI-systemen vereist dat deze drie elementen in balans worden gehouden, zodat AI-systemen per saldo positieve impact maakt zowel binnen als buiten de organisatie.

Maak de afweging tussen wat kan, wat mag en wat wenselijk is en breng daar balans in aan

De AI act: een risicobaseerde aanpak

De AI Act is een Europese wet die regels stelt voor de ontwikkeling en het gebruik van AI-systemen. Ook geeft de wet rechten aan burgers die in aanraking komen met AI-systemen. Het doel van de wet is dat AI-systemen die organisaties (binnen de EU) gebruiken, veilig zijn en fundamentele rechten respecteren (1).

De AI Act gaat uit van een risicobaseerde aanpak. Dit houdt in dat afhankelijk van het domein waar deze ingezet wordt, een bepaalde risicocategorie en daarbij behorende set aan regels gelden. In totaal kent de verordening drie risiconiveaus: onacceptabel, hoog en beperkt. Het uitgangspunt hierbij is dat hoe groter het risico, des te strenger de regels zijn die hiervoor gelden. De risicocategorie van een systeem hangt nauw samen met het toepassingsgebied en ziet er dus op toe dat er strengere regels zijn als de toepassing ervan kritischer is, ofwel meer impact heeft. Daarnaast gelden er specifieke vereisten voor leveranciers van generieke AI-systemen om ook oplossingen met een generieke toepassing te reguleren.

Systemen met een onacceptabel risico zijn door de EU vanaf februari 2025 niet toegestaan (1). Een onacceptabel risico kenmerkt zich door ethische onwenselijke activiteiten. Vaak zijn deze activiteiten al vanuit andere wetgeving verboden zoals etnische profilering.

In de beperkt-risicocategorie vallen AI-systemen die content genereren (generatieve AI), voor deze toepassingen geldt een transparantieverplichting.

Voor AI-systemen in 'hoog risico'-toepassingsgebieden gelden aanvullende verplichtingen, zoals het inrichten van een risico- en kwaliteitsmanagementsysteem in de organisatie van de Gebruiksverantwoordelijke (bijvoorbeeld

ISO42001), uitvoeren van een impact assessment op mensenrechten (bijvoorbeeld IAMA (2)), registreren in een Europees (algoritme)register en eigen administratie; implementeren van loggingsmogelijkheden en bewaartermijnen, inrichten van 'human oversight' en treffen van passende beveiligingsmaatregelen (3).

De risicobenadering uit de AI-act classificeert alleen AI-systemen in specifieke toepassingsgebieden als hoog risico. De bepalingen gelden dus niet voor algoritmen met hoge impact:

- ook buiten deze domeinen kunnen AI-systemen grote impact hebben op betrokkenen;
- ook algoritmen, die niet onder de definitie van AI vallen, kunnen impact hebben op betrokkenen.

De publicatiestandaard voor het algoritmeregister (4) van Nederlandse Overheden bevat een hulpmiddel voor de selectie voor publicatie van algoritmen. Het is te adviseren om ook voor de algoritmen met impact op betrokkenen (categorie B) eveneens de afweging te maken dit ook te integreren in je risico- en kwaliteitsmanagementsysteem om passende aanvullende maatregelen te nemen zoals het uitvoeren van een IAMA en/of adequate menselijke tussenkomst.

Risicoclassificatie en risico-assessments bij AI-systemen

Bij een risicoclassificatie voor een AI-systeem zul je dus al snel beginnen met het inschatten van de risicocategorie volgens de AI-Act en – afhankelijk van het beleid van je organisatie – ook concluderen of er sprake is van een impactvol algoritme. Het gaat hierbij dus vooral om hoe het algoritme wordt ingezet.

Het uitvoeren van risico-assessments voor AI verschilt op belangrijke punten van de standaard IT-risicoanalyses. Dit is te verklaren doordat traditionele IT-systemen anders werken dan AI-systemen. Traditionele IT-systemen volgen vaak vastgestelde parameters en deterministische processen, waar AI extra complexiteit en onvoorspelbaarheid introduceert doordat het juist ontwikkeld is om zelf patronen te herkennen. De parameters staan dus niet meer vast. Hierdoor gaan andere risico's een rol spelen. Denk hierbij aan:

Zelflerende en adaptieve systemen

Doordat AI-systemen zelf patronen herkennen en hierdoor de gevonden patronen dynamisch kunnen aanpassen op basis van nieuwe gegevens, kan het voorkomen dat deze

systemen bijvoorbeeld onbedoeld onjuiste of zelfs schadelijke patronen aanleren.

Uitlegbaarheid

Het is vaak moeilijk om precies te verklaren hoe een AI-systeem tot een bepaalde beslissing komt. In een traditionele IT-context is het over het algemeen eenvoudiger om beslissingen goed te herleiden. Voor elke context zal afgewogen moeten worden of beslissingen van het systeem wel voldoende uitlegbaar zijn.

Gevoeligheid voor bias (5)

AI leert patronen op basis van data. Hierdoor kan bias ontstaan als deze data (trainingsdata) niet representatief zijn of zelfs expliciete of impliciete vooroordelen bevatten. Dit kan leiden tot ongewenste effecten, zoals ongelijke behandeling.

Impact van autonome beslissingen

Wanneer de resultaten van een AI-systeem gebruikt worden om autonoom beslissingen te nemen, zoals het automatisch blokkeren van een IP-adres bij de detectie van verdachte activiteiten, kan het zijn dat er onterecht verdachte activiteiten zijn gedetecteerd. Dit komt doordat het een inschatting is en geen zekerheid. Een fout in de besluitvorming kan directe operationele gevolgen hebben. Er moet een afweging gemaakt worden wat enerzijds de foutmarge is, die geaccepteerd wordt, en wat anderzijds de gevolgen zijn van het missen van verdachte activiteiten. Om algoritmerisico's te identificeren en beheersen ontkomt je er niet aan om ook voor algoritmen risico- en kwaliteitsmanagement te borgen in de organisatie. Begin vandaag nog met de uitvoering en verbeter gaandeweg, in plaats van alles vooraf tot in detail uit te werken. Door te doen, ga je leren wat wel werkt en ook wat niet. Zo ervaar je samen wat nog verder, duidelijker of anders uitgewerkt moet worden om een 'AI-volwassen'-organisatie te worden. Wil je dit gestructureerd aanpakken en continue blijven verbeteren? Start met het inrichten en implementeren van een AI-governance.

AI-volwassenorganisatie met een AI-governance

AI-governance kun je zien als het totaalpakket van (beleids)afspraken, beleidskaders, processen, procedures, taken, rollen en verantwoordelijkheden; om enerzijds de risico's van AI-systemen te beheersen en anderzijds de kansen van AI maximaal en op verantwoorde en efficiënte



Onderdelen van AI-governance (6).

wijze te benutten. AI-governance is vaak een verlengde of onderdeel van een data- of IT-governance.

AI-governance bestaat dus uit een overzichtelijke set richtlijnen waarin jouw organisatie haar visie op AI, het beheersen van risico's en de verdeling van verantwoordelijkheden vastlegt. Denk hierbij aan een AI-visie, ethisch kader, risicomanagementprocessen en een duidelijke beschreven rolverdeling.

Door in een AI-governance vast te leggen wat afspraken zijn, processen en procedures te beschrijven en duidelijk toe te lichten wie wat moet doen hierin, ontstaat duidelijkheid over ieders rol om verantwoord met AI-systemen om te gaan. Ook beschrijf je tijdens het opstellen van een AI-governance wat er allemaal nodig is aan kennis, capaciteit en middelen om de ambities in de governance te realiseren. Dit faciliteert het organiseren van de benodigde capaciteit, middelen en budget voor succesvolle implementatie van verantwoorde AI-systemen.

Als organisatie bepaal je zelf de scope van deze governance. Bijvoorbeeld wanneer er voor informatiebeveiliging of privacy al veel risico's afgedekt worden. Of om als overheidsorganisatie in lijn te komen met de publicatiestandaard voor het algoritmeregister door te kiezen voor een bredere scope dan AI-systemen, een algoritmegovernance.

Is het mogelijk om als organisatie verantwoord om te gaan met AI?

Het belangrijkste punt hierbij is dat de afweging tussen wat kan, wat mag en wat wenselijk is in balans gebracht kan worden. Bijvoorbeeld doordat de governance rondom AI ingericht is en er gedegen risicomanagement plaatsvindt. Zo zorg je als organisatie dat je in control bent en blijft van AI-systemen in gebruik.

Referenties

- (1) <https://www.linkedin.com/pulse/power-trio-how-ai-quantum-computing-cybersecurity-our-lokhandwala-hlxf>
- (2) <https://www.ibm.com/case-studies/cargills-bank-ltd>
- (3) https://www.splunk.com/en_us/customers/success-stories/tide.html
- (4) <https://www.trellix.com/assets/case-studies/trellix-election-casestudy.pdf>
- (5) <https://www.wsj.com/articles/fraudsters-use-ai-to-mimic-ceos-voice-in-unusual-cybercrime-case-11567157402>
- (6) <https://spectrum.ieee.org/real-time-deepfakes>
- (7) <https://drlee.io/ai-security-model-hacking-with-model-inversion-attacks-techniques-examples-and-real-world-a23b5fff272a>
- (8) <https://neptune.ai/blog/adversarial-machine-learning-defense-strategies>

Auteur: Inge Wetzter is werkzaam bij Secura als social psychologist cybersecurity & compliance. Zij is gepromoveerd in de sociale psychologie en is bereikbaar via e-mail: inge@secura.com



De psychologie over een effectief mensprogramma in cybersecurity

Baat het niet dan schaadt het wel

Deel 1 van tweeluik

Cybersecurityprogramma's focussen steeds meer op gedrag, maar zonder doordachte aanpak werkt dit vaak averechts. Veelgemaakte fouten zijn: teveel aandacht voor wat al goed gaat, focus op verkeerde gedragsfactoren zoals kennis of motivatie, en mismatches met het kennisniveau. Dit leidt tot frustratie, weerstand en onveilig gedrag, ondanks goede intenties. Door eerst te meten wat medewerkers weten en de echte knelpunten te identificeren, kunnen organisaties gerichte interventies inzetten en duurzame gedragsverandering bereiken zonder weerstand.

DORA bestaat uit vijf pijlers, waarvan het beheer van ICT-risico's van derde aanbieders er één is. Dit laat het belang van het managen van uitbestedingsrisico's zien. De specifieke eisen zijn vastgelegd in artikelen 28 t/m 30. Artikelen 31 t/m 44 bevatten de regels rondom het overzichtskader, maar die zijn op dit moment voor financiële instellingen minder relevant.

Je opent je inbox en ziet weer een nieuwe cybersecuritytraining verschijnen. Het is al de derde dit jaar en je hebt de vorige pas net afgerond. Wederom krijg je te horen hoe belangrijk sterke wachtwoorden zijn en dat je alert moet blijven op phishingmails. Je weet het allemaal wel inmiddels. Bij de koffieautomaat hangt weer zo'n poster: "Vergrendel je scherm als je wegloopt!" En dan zijn er nog de gele kaarten: de zoveelste waarschuwing, de zoveelste reminder—het begint zijn effect te verliezen. Sterker nog, het begint langzaam weerstand op te roepen.

De afgelopen jaren is er binnen organisaties steeds meer aandacht gekomen voor de menselijke factor in cybersecurity.

Waar cybersecurity voorheen voornamelijk werd gezien als een technische kwestie, gericht op firewalls, antivirussoftware en encryptie, beseffen steeds meer bedrijven dat het gedrag van medewerkers minstens zo belangrijk is om cyberdreigingen tegen te gaan. Cyberaanvallen worden steeds geraffineerder, maar tegelijkertijd blijkt uit onderzoek dat een groot deel van de beveiligingsincidenten nog steeds voortkomt uit menselijk handelen: onbewuste fouten, slordigheden of simpelweg het niet volgen van protocollen.

Dit inzicht heeft ertoe geleid dat organisaties massaal zijn gaan inzetten op het vergroten van cybersecurity-awareness onder hun personeel. Via trainingen, e-learning, posters en workshops worden medewerkers bewust gemaakt van de risico's en verantwoordelijkheden. Denk aan waarschuwingen over phishingmails, het gebruik van sterke wachtwoorden en het belang van het vergrendelen van je computerscherm als je wegloopt. Deze focus op bewustwording – of awareness – is een belangrijke eerste stap in het erkennen van de rol van de mens in cybersecurity.

Meer aandacht voor gedrag

De afgelopen jaren heeft een belangrijke verschuiving plaatsgevonden: veel organisaties hebben ingezien dat het niet enkel draait om bewustwording, maar juist om gedragsverandering. De psychologie leert ons dat gedrag wordt bepaald door meer dan alleen bewustzijn. Naast bewustzijn spelen motivatie (wil iemand het wel?) en gelegenheid (wordt iemand in staat gesteld om het te doen?) een cruciale rol. Weten wat veilig gedrag is, betekent niet automatisch dat iemand het ook vertoont. Mensen moeten zowel kennis hebben als gemotiveerd zijn én de mogelijkheid hebben om het gedrag daadwerkelijk te vertonen.

Deze inzichten komen voort uit onderzoek naar de menselijke kant van cybersecurity. Daarbij werd duidelijk dat gedrag slechts deels afhankelijk is van kennis. Zonder aandacht voor motivatie en gelegenheid blijft gedragsverandering uit ((1), (2), (3), (4)).

Veel organisaties hebben inmiddels begrepen dat het niet genoeg is dat medewerkers weten wat veilig gedrag is; ze moeten dit gedrag ook daadwerkelijk vertonen. Dit heeft geleid tot een toenemende focus op gedragsveranderingsprogramma's: posters ophangen in kantoren, medewerkers aanmoedigen elkaar aan te spreken op onveilig gedrag en het uitdelen van gele kaarten bij overtredingen.

Het idee achter deze programma's is goed: medewerkers actief betrekken bij cybersecurity en hen helpen om cyberveilig gedrag te omarmen. Maar hier schuilt ook een risico. In veel gevallen worden deze gedragsprogramma's ingevoerd zonder dat er grondig wordt onderzocht hoe medewerkers ze ervaren. De druk om snel iets te doen—om iets uit te rollen—kan ertoe leiden dat organisaties maar wat proberen, in de hoop dat er iets blijft hangen. In plaats van medewerkers te helpen hun gedrag structureel te veranderen, kunnen slecht doordachte gedragsprogramma's juist negatieve effecten hebben.

De risico's van ongefundeerde initiatieven

Een overdaad aan losse interventies – van herhaalde trainingssessies en postercampagnes tot het uitdelen van gele kaarten – kan ertoe leiden dat medewerkers vermoeid raken en afhaken. In plaats van cyberveiligheid als iets positiefs en belangrijks te zien, ontstaat er irritatie of zelfs cynisme. Wanneer gedragsprogramma's niet goed aansluiten bij de werkpraktijk of de motivatie

van medewerkers, kunnen ze meer kwaad dan goed doen. Dit betekent dat organisaties zorgvuldig moeten nadenken over hun aanpak van gedragsverandering.

De sleutel ligt in een doordachte, op maat gemaakte benadering, die niet alleen kijkt naar kennis en bewustwording, maar ook naar de motivatie, gelegenheid en ondersteuning die nodig zijn om gedragsverandering daadwerkelijk te laten plaatsvinden. Dit artikel beschrijft hoe gedragsprogramma's effectief kunnen worden ingezet om medewerkers te ondersteunen, zonder hen te overspoelen met ondoordachte campagnes die uiteindelijk tot weerstand leiden.

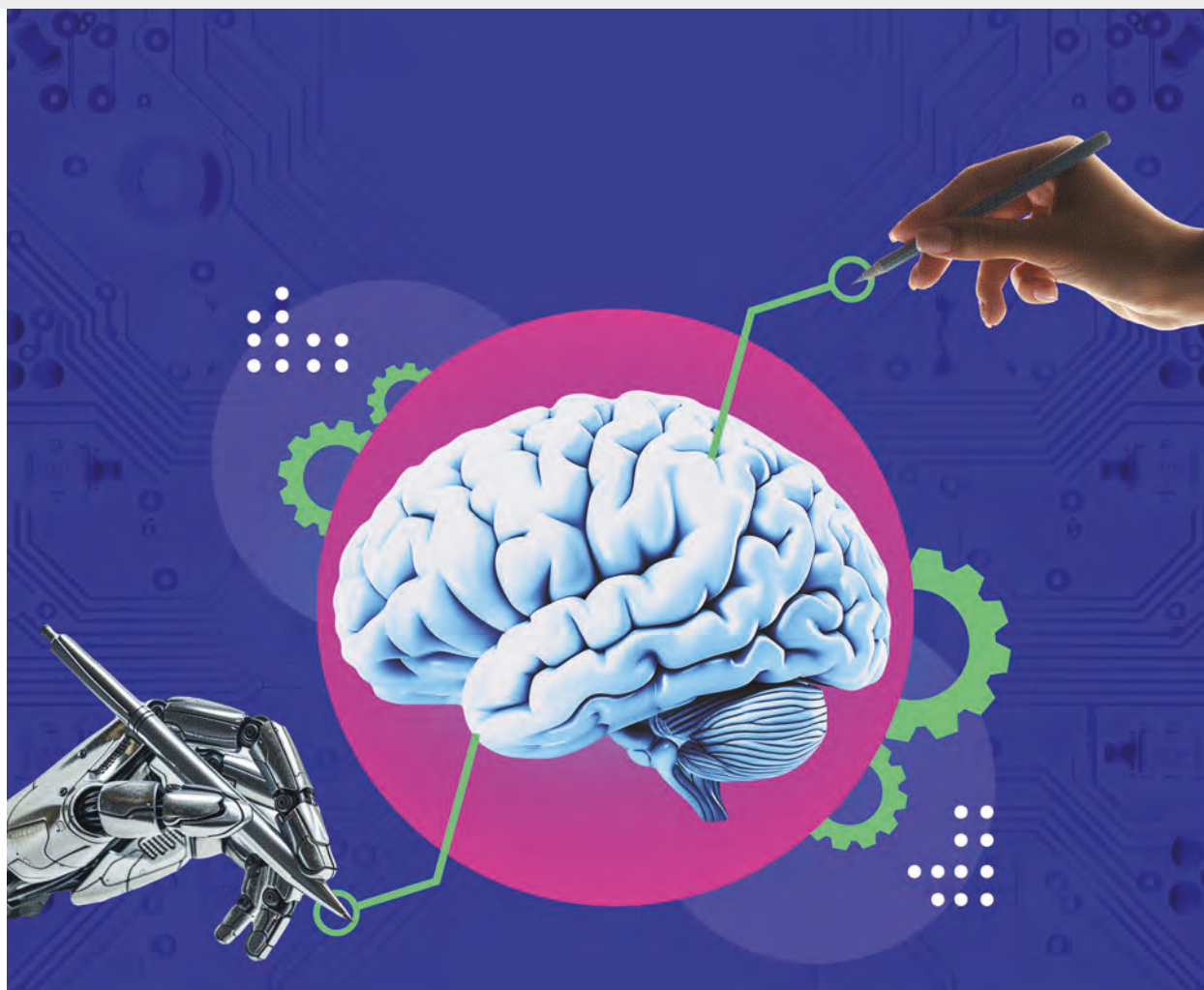
De meest gemaakte fouten in gedragsprogramma's

Wanneer een gedragsprogramma blijft hameren op onderwerpen die medewerkers al beheersen, kan dit snel leiden tot frustratie en weerstand. Mensen worden moe van het steeds opnieuw horen van dezelfde boodschap, zeker als ze al weten wat er van hen verwacht wordt. Dit kan resulteren in wat we security fatigue noemen: een toestand waarin medewerkers verzadigd raken door herhaling en hun motivatie om nog mee te doen afneemt. Ze haken af, niet omdat ze de risico's niet serieus nemen, maar omdat ze het gevoel krijgen dat de programma's weinig toevoegen. Dit maakt het lastig om hen later nog te bereiken met nieuwe, mogelijk wél relevante informatie.

De oplossing begint met inzicht. Voordat je een gedragsprogramma uitrolt, moet je eerst vaststellen welke onderwerpen al goed gaan en waar nog werk aan de winkel is. Dit betekent niet zomaar aannemen of gokken, maar echt meten. Alleen door te weten wat medewerkers al goed doen en waar de knelpunten liggen, kun je een programma opzetten dat zich richt op wat er écht toe doet. Zo voorkom je dat medewerkers verzanden in herhaling en frustratie, en zorg je ervoor dat hun aandacht gericht blijft op de onderwerpen die daadwerkelijk verbetering nodig hebben.

Aandacht voor de verkeerde factor van gedrag

Een tweede veelgemaakte fout in gedragsprogramma's is dat de focus ligt op de verkeerde factor van gedrag. Vaak wordt kennis eindeloos herhaald, terwijl het echte probleem in de gelegenheid zit. Medewerkers weten wel wat ze moeten doen en



willen het ook wel, maar worden belemmerd door praktische obstakels. Dit leidt tot frustratie: als je steeds hoort wat je al weet, zonder dat de juiste omstandigheden worden gecreëerd om het gedrag daadwerkelijk te vertonen, voelt dat als zinloze herhaling. Hierdoor ontstaan workarounds: medewerkers zoeken manieren om hun werk toch gedaan te krijgen, vaak op manieren die de veiligheid ondermijnen. Zolang de barrières in de gelegenheid niet worden weggenomen, blijft het gewenste gedrag uit en groeit de weerstand tegen verdere interventies.

Het kan natuurlijk ook andersom: dat er juist aandacht voor kennis nodig is, maar dat de nadruk ligt op het verhogen van motivatie. Medewerkers willen wel veilig handelen, maar weten niet precies hoe. Neem bijvoorbeeld het kiezen van een sterk wachtwoord: wat in 2025 nog als veilig wordt beschouwd, verandert snel. Als deze kennis niet up-to-date is, heeft het weinig zin om medewerkers te blijven motiveren zonder hen de juiste informatie te geven. Dit leidt niet alleen tot verwarring, maar ook tot onveilig gedrag, simpelweg omdat mensen niet weten wat van hen

wordt verwacht. Het resultaat is dat, ondanks alle inspanningen om motivatie te verhogen, het gedrag niet verandert omdat de basiskennis ontbreekt.

Om effectief gedrag te veranderen, moet je per onderwerp helder krijgen welke factor ontbreekt: kennis, motivatie of gelegenheid. Dit betekent dat je niet zomaar kunt aannemen wat het probleem is, maar dit daadwerkelijk moet meten. Als het bijvoorbeeld gaat om veilig wachtwoordgebruik, moet je weten of medewerkers niet begrijpen wat een sterk wachtwoord is, of dat ze het wel weten, maar niet gemotiveerd zijn om het toe te passen, of dat ze worden belemmerd door praktische obstakels zoals ingewikkelde procedures. Pas als je precies weet waar het knelt, kun je gericht ingrijpen. Zo voorkom je dat je blijft hameren op kennis terwijl de barrières in de faciliteiten liggen, of dat je motivatie probeert te verhogen terwijl mensen simpelweg niet weten wat ze moeten doen. Alleen met een gerichte aanpak kun je echte gedragsverandering realiseren.

Effectieve gedragsprogramma's in cybersecurity vragen om meer dan alleen goede bedoelingen

Aandacht op het verkeerde kennisniveau

Een derde veelgemaakte fout in gedragsprogramma's is dat de inhoud niet is afgestemd op het juiste kennisniveau van de medewerkers. Soms zijn de e-learnings of trainingen te simpel, waardoor mensen het gevoel krijgen niet serieus genomen te worden. Dit kan leiden tot irritatie en demotivatie, omdat ze het idee krijgen dat hun kennis wordt onderschat of dat er niet genoeg moeite wordt gedaan om de inhoud op maat aan te bieden, terwijl er wel tijd wordt gevraagd om een training te volgen bovenop een toch al vol takenpakket. Aan de andere kant kunnen de programma's ook te complex zijn, waardoor medewerkers overweldigd raken. Dit kan resulteren in een gevoel van machteloosheid of learned helplessness, waarbij ze denken dat ze het toch niet kunnen begrijpen of toepassen. In beide gevallen raken medewerkers gedemotiveerd en neemt hun betrokkenheid af, wat de effectiviteit van het programma ondermijnt.

Om dit te voorkomen, is het essentieel om eerst het juiste kennisniveau van de medewerkers te meten voordat je een training of e-learning aanbiedt. Door te weten waar je medewerkers staan, kun je de inhoud van het programma beter afstemmen op hun behoeften. Dit betekent dat je voorkomt dat trainingen te eenvoudig zijn voor mensen die al veel weten, en niet te complex voor degenen die nog basiskennis nodig hebben. Een goed afgestemd programma zorgt ervoor dat iedereen zich serieus genomen voelt en op het juiste niveau wordt uitgedaagd. Zo blijven medewerkers betrokken en voorkom je dat ze afhaken door frustratie of overweldiging.

Belangrijke lessen

Effectieve gedragsprogramma's in cybersecurity vragen om meer dan alleen goede bedoelingen. Ze vereisen een zorgvuldige aanpak waarbij organisaties niet alleen inzicht krijgen in de huidige kennis en het gedrag van hun medewerkers, maar ook begrijpen welke factoren echte verandering tegenhouden. Door te meten en te richten op wat werkelijk nodig is of dat nu

kennis, motivatie of gelegenheid betreft—kunnen organisaties programma's ontwikkelen die niet alleen effectief zijn, maar ook weerstand en frustratie voorkomen. Het doel is niet alleen om medewerkers bewuster te maken, maar hen ook daadwerkelijk in staat te stellen en te motiveren om veilig gedrag te vertonen, waardoor de organisatie als geheel sterker en veiliger wordt. In deel 2 van dit tweeluik zullen cijfers gepresenteerd worden van een grote benchmarkstudie. Deze resultaten laten zien waar organisaties gemiddeld genomen staan: welke onderwerpen blijven achter in gedrag en waar ontbreekt het aan kennis, motivatie of gelegenheid? Door de resultaten van een grote steekproef over verschillende organisaties te analyseren, krijgen we inzicht in de startpunten (gemiddeld gezien) van een effectief programma. Want inmiddels weten we: als gedragsprogramma's slecht worden ontworpen, kunnen de gevolgen ernstig zijn. Medewerkers kunnen gefrustreerd raken door herhaling van overbodige informatie, wat leidt tot demotivatie en cynisme. Onjuiste focus op kennis of motivatie kan ervoor zorgen dat belangrijke praktische barrières niet worden weggenomen, waardoor medewerkers workarounds gaan zoeken die de veiligheid ondermijnen. In het ergste geval ontstaat er een gevoel van machteloosheid, waarbij medewerkers afhaken en zelfs bewust onveilig gedrag vertonen. Dit benadrukt hoe cruciaal het is om zorgvuldig te meten en de juiste interventies in te zetten. Want wat betreft 'awareness' in cybersecurity geldt: baat het niet, dan schaadt het wel.

Referenties

- (1) Wetzer, I. M. (2017a). Voorbij awareness; grip op cyberveilig gedrag. *InformatieBeveiliging*, 17 (IB.3), 24-26.
- (2) Wetzer, I. M. (2017b). Cyberveilig gedrag: Meer dan alleen het locken van je beeldscherm. *InformatieBeveiliging*, 17 (IB.6), 4-7.
- (3) Wetzer, I. M. (2018). Cyberveilig gedrag: Waarom dóen we het nou niet? *InformatieBeveiliging*, 18 (IB.1), 12-15.
- (4) Wetzer, I. M. (2021). Het begint met bewustwording. Hoe ver zijn we daar inmiddels mee? *Informatie Beveiliging*, 6, 26-29.

Lex Borger is security consultant bij Tesorion en oud-hoofdredacteur van IB Magazine. Lex is bereikbaar via lex.borger@tesorion.nl



AI in control

Aan het begin van dit nieuwe decennium, in deze eerste editie van 2030, wil ik mijn enthousiasme delen over de nieuwste beta-release van het ISMS-tool Trusted in Control 3.0. Wie had gedacht dat ik, na al die jaren als CISO, nog zoveel plezier zou beleven aan het dagelijks managen van informatiebeveiligingsuitdagingen? Met deze tool voel ik me eindelijk écht in control. Het biedt niet alleen direct inzicht en overzicht, maar geeft ook weloverwogen advies en de mogelijkheid om tot op het kleinste detail door te dringen. Dat geeft vertrouwen. Ja, je kunt het niet helpen om dan even te glunderen...

Gisteren liep ik nog risicoanalyses door voor twee afdelingen. Het ISMS had alles al voorbereid: gebruikersgedrag en applicatiegebruik waren geanalyseerd, threat intelligence was vergaard, en de business impact werd gemodelleerd tegen de meest recente geopolitieke ontwikkelingen. Het resultaat? Een heldere samenvatting van aanpassingen en aanbevelingen, die ik vervolgens samen met de verantwoordelijke managers doornam. Ze gingen direct akkoord. De behandelplannen? Automatisch ingediend bij portfolio management, inclusief prioritering ten opzichte van andere lopende initiatieven.

Na de koffie nam ik nog even het door het ISMS gegenereerde rapport door over de naleving van de ISO 27001-structuur. Op enkele kleine verbeterpunten na – die overigens al automatisch verwerkt zijn in de behandelplannen – zag alles er piekfijn uit. Drie principiële punten in het autorisatiebeleid vroegen om een kleine aanpassing vanwege de nieuwste technische updates in Defender voor Windows 13. Geen stress: deze punten staan ook automatisch op de agenda voor de volgende stuurgroepvergadering.

Onze leveranciers blijven trouw hun data aanleveren, de meeste zelfs on-demand via de nieuwste IIEP*-versie 2.1. Een paar onderleveranciers hebben hun processen nog niet op orde, maar omdat hun bijdrage onder onze risicodrempel valt, hoeven we ons daar niet meer druk om te maken. Het systeem regelt het.

Een ander indrukwekkend aspect van Trusted in Control 3.0 zijn de integraties voor het automatisch scannen en beoordelen van kwetsbaarheden. Dagelijks monitort het systeem onze IT-omgeving op bekende en opkomende dreigingen, inclusief softwarekwetsbaarheden, configuratiefouten en ongepatchte systemen.

Wat deze tool uniek maakt, is de intelligentie achter de kwetsbaarheidenscan. Het gaat niet alleen om detectie, maar ook om context. Zodra een kwetsbaarheid wordt ontdekt, beoordeelt het ISMS direct de potentiële impact op onze specifieke infrastructuur en prioriteert de aanpak op basis van het risico. Wat geautomatiseerd in gang gezet kan worden, wordt meteen afgehandeld.

Na de lunch wierp ik een blik op het ISMS-dashboard. Alle KPI's staan netjes op groen. Afwijkingen? Die worden al gecorrigeerd voordat iemand ze überhaupt kan waarnemen. Volgende week staat de certificatie-audit op de planning, en alle seinen staan op groen. Het is bijna jammer dat ik niet nú al met een druk op de knop alle auditrelevante gegevens direct naar de auditor mag sturen, dat komt volgende week. Het stage 1-rapport wordt dan automatisch gegenereerd, compleet met een voorstel voor het stage 2-rapport. Op het eind van de dag hebben we het in handen.

De stage 2 audit vereist nog een controlebezoek, omdat de certificeringsinstanties dit formeel voorschrijven. Maar wie weet hoe lang nog? PCI-DSS-audits verlopen inmiddels volledig automatisch. Dat ik dit nog mag meemaken!

Het is een koude maar zonnige winterdag. Alles loopt gesmeerd; ik besluit een frisse neus te halen met een boswandeling. Terwijl ik nadenk over nieuwe security-initiatieven en een paar afleveringen van mijn favoriete podcast luister, realiseer ik me dat dit ISMS precies biedt wat technologie zou moeten doen: rust creëren, ruimte maken voor visie.

En dan... schrik ik wakker. Het is 2025. Een paar zweetdruppels lopen langs mijn voorhoofd. Ik kijk op de klok. Het is bijna tijd om op te staan en weer aan de slag te gaan. Welke worstelingen met het ISMS-tool wachten er vandaag op me?

*IIEP = ISMS Information Exchange Protocol

BLOG

Beter spreken in het openbaar

Minstens één van deze tips, ideeën en technieken over spreken in het openbaar zal ooit je leven veranderen. Het verschil tussen een afwijzing of een geslaagde sollicitatie. Een presentatie voor collega's zonder of met applaus na afloop. Iemand die na je verhaal komt zeggen dat zij nu heel anders tegen het behandelde probleem aankijkt. Deze tips komen uit een presentatie over presentaties die Patrick Henry Winston jaarlijks in januari gaf. Hij was Ford Professor Artificial Intelligence and Computer Science aan het Massachusetts Institute of Technology en overleed in 2019. Van zijn presentatie in 2018 is een video van een uur beschikbaar (1).



Je carrière wordt grotendeels bepaald door hoe goed je spreekt, hoe goed je schrijft en de kwaliteit van je ideeën. In die volgorde en in afnemend belang, volgens Winston. Officierien komen volgens het militaire recht voor de krijgsraad als ze hun soldaten zonder wapens het slagveld opsturen. Ook studenten moeten van hun docenten technieken meekrijgen om hun ideeën te kunnen 'verkopen'. Zonder die technieken is het als inhoudelijk deskundige security officer alsof je in je broek plast, terwijl je een donker pak draagt. Het geeft jou een warm gevoel, maar niemand ziet er iets van.

Winston schijft aan het begin van de presentatie met de hand deze formule op een enorm schoolbord:

$$I = f(K, P, T)$$

Je Impact (I) is een functie (f) - we zijn immers op MIT - van je Kennis (K) over spreken, hoeveel je oefent in de Praktijk (P) en je aangeboren Talent (T). Ook hier in volgorde van afnemend belang: genetisch talent voor spreken, als het al bestaat, speelt amper een rol.

De adviezen van Winston richten zich op de kennis over spreken. Dit is de gemakkelijkste manier om de grootste toename van impact te behalen. Spreken in het openbaar doe je vaak. Zakelijk: bij het geven van een security awareness presentatie of bij een sollicitatie. Privé: bij speeches voor familieleden of als 'best (wo)man' op een huwelijk of bij een presentatie aan bejaarden over het tractorgebruik in de Achterhoek.

Een presentatie heeft een begin, een middenstuk en een eind. Vaak zijn er vragen vanuit het publiek. Er worden vrijwel altijd hulpmiddelen zoals PowerPoint gebruikt en Winston adviseert ook verhalen en attributen te gebruiken. Al deze onderwerpen komen hierna aan bod.

Het begin van de presentatie

Winston vraagt iedereen om laptops en telefoons weg te doen. Mensen hebben maar één taalverwerker in hun hersens en die kan maar één ding tegelijk. Hij wil dat ze naar hem luisteren en niet hun mail bijwerken. Begin je presentatie nooit met een grap, want je publiek kent je niet en moet wennen aan je specifieke manier van praten en mimiek. Daardoor mislukt een grap aan het begin vaak. Geef de luisteraars wel een 'menu': vertel waar je het over gaat hebben en in welke volgorde. Dit geeft hen een rode draad. Je kunt heel goed starten met een belofte: je gaat je publiek iets positiefs, nuttigs of tenminste interessants leveren. Doe dit op een enthousiaste manier, zodat ze het gevoel hebben dat dit de belangrijkste en interessantste speech ooit wordt en dat ze heel

veel gaan leren. Dus niet: 'Ik ga iets over SPAM vertellen omdat ze mij dat vanmorgen hebben gevraagd...'

Het middenstuk: de 'body'

Winston behandelt vier aanpakken voor het middenstuk:

1. Cycling - 'rondfietsen'. Vertel over je idee, beschrijf dan de details van je idee en kom nog een derde keer terug op de kern van je idee. Zodat de mensen die de draad kwijtraken, toch denken dat ze tenminste de hoofdlijn hebben begrepen. Dus: in het kort, in detail en in samenvatting. Of in termen van artificial intelligence: vertel ze het schema te laden, vul dan de details in, laat ze daarna weten wat geïndexeerd moet worden voor de toekomst.
2. Verbal punctuation - 'hapklare brokken'. Elk moment van je presentatie is 20% van de luisteraars (gelukkig niet steeds dezelfde!) 'even afgeleid'. Hak je presentatie op in delen en vertel je publiek tussen die delen door waar je bent. Dus: 'Ik heb het gehad over cycling, nu ga ik het hebben over verbal punctuation en daarna heb ik nog een verrassing voor jullie'. Door de opdeling kunnen luisteraars de rode draad weer oppakken.
3. Fence your idea (or Near Miss) - 'wat is het niet?'. Ideeën lijken op abstract niveau bekeken allemaal op elkaar. Leg daarom uit welke ideeën wel veel lijken op jouw idee, maar toch anders zijn. Bijvoorbeeld: Het algoritme van Piet Puk levert ook een antwoord, maar zijn algoritme is exponentieel en dat van mij is lineair, en ik ben daarom klaar vóórdat het zonnestelsel explodeert en de aarde smelt.
4. Ask Rhetorical Questions. Maak ze niet te gemakkelijk, want dan durft niemand iets te zeggen. Maak ze ook niet te moeilijk. Wacht ongeveer zes seconden op een antwoord en ga dan verder met je verhaal.

Tijd en plaats

11:00 uur is een prachtige tijd voor een presentatie. Iedereen - zelfs MIT-studenten -, is inmiddels wakker, er is nog bijna niemand die alweer gaat slapen. Het is niet net na een (zware) maaltijd, je zit niet tussen je publiek en de borrel na afloop.

Mensen hebben de neiging om in een schemerige of donkere kamer hun ogen te sluiten. Dan is het moeilijk om jouw schoolbord, whiteboard, flipover of PowerPoint-scherm te lezen. Zorg dus dat de ruimte steeds goed verlicht is. Winston stelt: als je een bank gaat beroven, ga je de zaak toch ook eerst 'afleggen' (= verkennen voor criminele doeleinden (2))? Zo vermijd je onaangename verrassingen voor jou als spreker.

Wanneer meer dan de helft van de stoelen in de zaal gevuld is, denken de aanwezige personen niet dat er elders iets gebeurt dat

Je publiek onthoudt de geschreven informatie op je slides beter dan je gesproken verhaal. Als je dat echt wilt, kun je beter een artikel schrijven en dat verspreiden

veel interessanter is. Stem dus de zaalomvang af op de verwachte hoeveelheid luisteraars.

Hulpmiddelen

Met een schoolboard of whiteboard kun je grafische elementen gebruiken. Iets omcirkelen, een pijl van het ene naar het andere samenhangende onderwerp, of juist een kruis zetten door een slecht idee. Het tempo van schrijven op een bord komt overeen met de snelheid waarmee mensen nieuwe informatie kunnen verwerken. Dit tempo helpt je ook om niet te snel te spreken tijdens je presentatie. Er zijn weinig presentatoren die uit zichzelf te langzaam praten. Bovendien heb je iets om met je handen naar te wijzen tijdens je presentatie. Winston legt uit dat in sommige culturen het als heel negatief wordt ervaren wanneer je jouw handen in de zakken of op je rug houdt: misschien verberg je wel een wapen! Door op het bord te schrijven, herinneren je luisteraars zich dat ze dat vroeger ook weleens deden: met de hand schrijven en nadenken over een onderwerp. Dit geeft meer verbinding met je publiek.

Gebruik geen laserpointer. Om het rode puntje ergens op te mikken, moet je met je rug of in elk geval je achterhoofd naar het publiek staan. Daardoor verlies je contact. Als je zenuwachtig bent, ziet iedereen dat rode puntje trillen. Gebruik ook geen houten aanwijspijp of zo'n uitschuifbare metalen. In zijn presentatie laat Winston een aanwijspijp zien, om die vervolgens doormidden te breken. Als je iets kunt laten zien met een tastbaar voorwerp, werkt dat altijd beter dan een foto of tekst op een slide. Een voorwerp dat je nog niet gebruikt, geeft je luisteraars verder het gevoel 'slim' te zijn omdat ze alvast zelf kunnen bedenken waarom je dat ding hebt meegebracht.

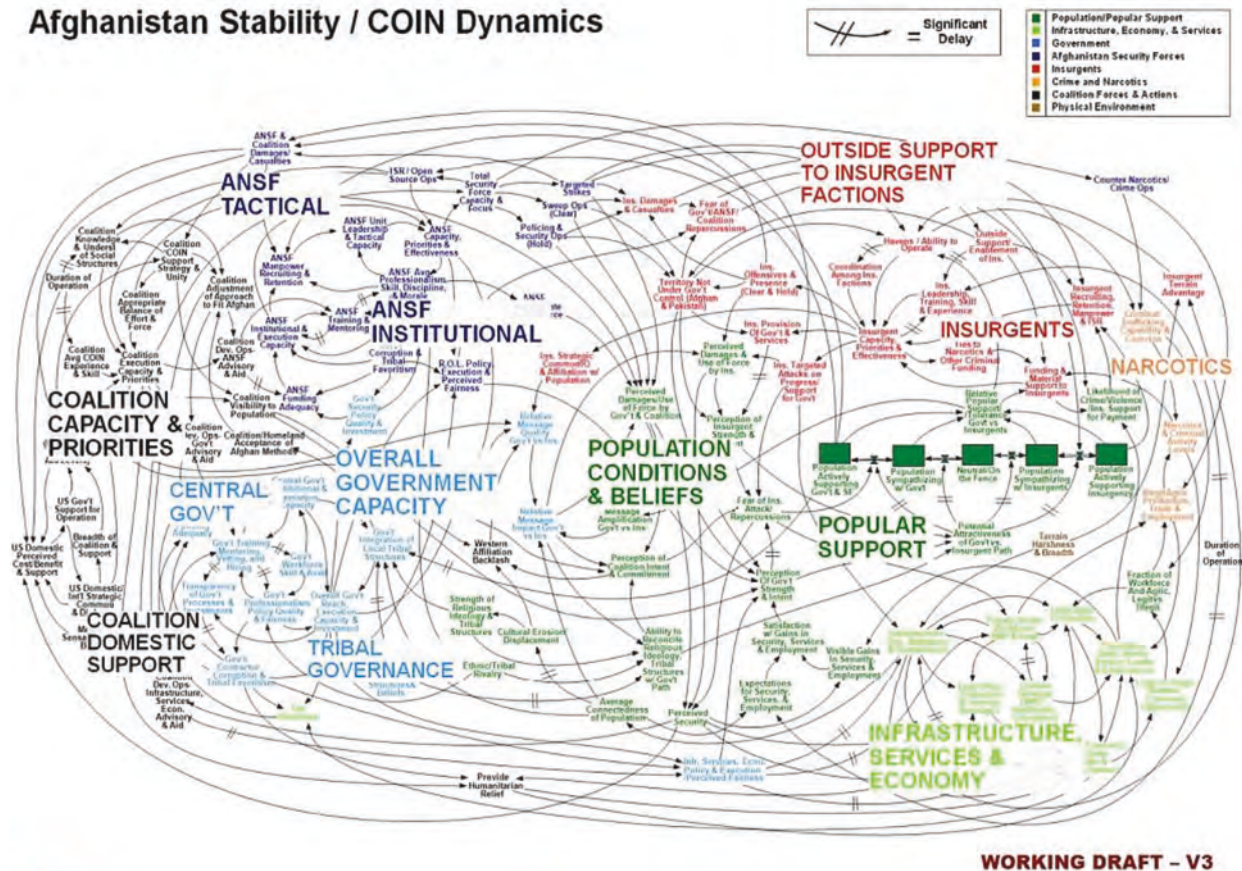
Slides

Gebruik het bord of flipover om iets uit te leggen of te onderwijzen. Gebruik slides om iets te laten zien of om mensen te overtuigen dat jij gelijk hebt met je aanpak of idee. Gebruik op slides nooit een lettertype kleiner dan 40 punten. Anders ga je teveel tekst op de slide zetten. Op reis vroeg iemand ooit aan Winston om commentaar te geven op zijn slides. 'Je hebt er te veel en er staat te veel op' zei hij, voordat hij de slides had gezien. Je presentatie afdrukken (één slide per A4) en dan alles op een tafel leggen, geeft je snel een idee of jij ook 'te veel' slides hebt en of ze 'te vol' zijn. Lees je slides niet hardop voor. Zeker als je over artificial intelligence spreekt, zoals Winston, of information security, zoals wij, mag je ervan uitgaan dat iedereen in de zaal zelf kan lezen. Blijf in de buurt van het scherm staan, zodat je luisteraars niet hun hoofd heen en weer hoeven te draaien zoals bij een tenniswedstrijd.

Gebruik één betekenisvol plaatje per slide. Een paperclip-manneke in de bijgeleverde 'art' van PowerPoint is niet betekenisvol. Zet er een pijl bij als een deel van het plaatje of schema heel belangrijk is. Nummer de pijlen als je er meer hebt. Zet er niet TIEN pijlen op! Een hapax legomenon is een woord dat slechts één keer in een boek of tekst voorkomt. Taalkundigen kunnen dus niet uit verschillende contexten afleiden wat het betekent. In je presentatie mag je één 'HL' gebruiken. Zoals een ingewikkeld schema dat niemand kan begrijpen, maar waardoor iedereen snapt dat het getoonde onderwerp heel ingewikkeld is.

Je publiek onthoudt de geschreven informatie op je slides beter dan je gesproken verhaal. Als je dat echt wilt, kun je beter een artikel schrijven en dat verspreiden.

Afghanistan Stability / COIN Dynamics



WORKING DRAFT - V3

PA Consulting
Group
© PA Knowledge Limited 2009

Page 22

Dit is het einde

Herinner je publiek aan de belofte die je aan het begin hebt gedaan (zie mijn eerste zin), en toon samenvattend aan dat je deze hebt waargemaakt. Eindig je presentatie gerust met een grap, zodat je publiek het gevoel heeft dat je verhaal de hele tijd door plezierig was! Vraag je publiek om je vragen te stellen. Niet met een slide met daarop 'vragen?', die dan twintig minuten in beeld blijft. Daar kun je al die tijd beter een foto van jezelf, je naam, eventuele certificeringen en je mailadres laten zien. Bedenk een vraag die je aan jezelf kunt stellen om de zaal op gang te krijgen. Herhaal de vraag uit de zaal zodat iedereen hem kan verstaan. Probeer een gesprek te voeren met de vraagsteller in plaats van opnieuw een lezing te geven. Kap een zogenaamde

'doorvragers' resoluut af en behoud zelf de regie over je presentatie. Bedank aan het eind je publiek NIET, en zeker niet 'voor hun tijd'. Het klinkt dan alsof ze uit beleefdheid jou een plezier hebben gedaan door te blijven zitten, maar dat ze liever iets heel anders hadden gedaan. Wat je aan het eind wél heel goed kunt doen, is het publiek groeten en een compliment geven. Zoals: 'U was een geweldig publiek. Ik hoop dat u veel heeft geleerd over het geven van een goede presentatie.'

Referenties

- (1) Video: <https://ocw.mit.edu/courses/res-tll-005-how-to-speak-january-iap-2018/>
- (2) (Amsterdams) iemand in de gaten houden, bespioneren



Achter Het Nieuws

In deze rubriek geven enkele IB-redacteuren in een kort stukje hun reactie op recente nieuwsitems over informatiebeveiliging. Dit zijn persoonlijke reacties van de auteurs en deze geven niet noodzakelijkerwijs het officiële standpunt weer van hun werkgever of van PvIB. Vragen en/of opmerkingen kun je sturen naar ibmagazine@pvib.nl.



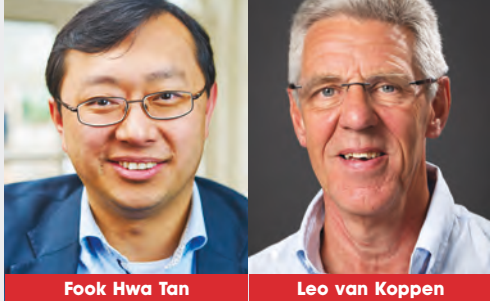
Cyber Resilience Act: tijd voor verantwoorde softwarekwaliteit

Het onderwerp van AHN in dit nieuwe jaar valt voor velen van u naar alle waarschijnlijkheid in de categorie 'klein nieuws'. Voor mij persoonlijk (Leo van Koppen) is het groot nieuws! Al zo'n twintig jaar verbaas ik me over, nee sterker nog, kijk ik met afschuw naar de wijze waarop softwareproducenten omgaan met hun verantwoordelijkheid. Hét nieuws van deze week gaat daar, naar ik hoop en ook wel verwacht, eindelijk verandering in brengen. De Cyber Resilience Act (CRA), die op 11 december 2024 van kracht is geworden, markeert een mijlpaal in softwareverantwoordelijkheid binnen de EU. Eindelijk gaan er eisen gesteld worden aan software, apps, hardware, IoT en digitale producten.

De wet is al aangekondigd in najaar 2022 en het heeft al twee jaar geduurd voordat er enige overeenstemming is gevonden in de uiteindelijke uitvoering. Met de stapsgewijze intrede hebben producenten de tijd om aan de eisen te gaan voldoen, ze hebben de tijd om de wettelijke eisen te vertalen naar de wijziging/invoering van de interne processen. Belangrijke data zijn: 11-12-2024 inwerkingtreding, 11-12-2026 meldplicht in werking, 11-12-2027 wetgeving volledig van kracht.

Belangrijkste elementen in de wetgeving:

1. In de ontwerp- en ontwikkelingsfase dienen security-eisen te worden meegenomen
2. Beperking van het aantal aanvalsoppervlakken zoals interfaces
3. Een risicobeoordeling dient onderdeel te zijn van het ontwikkelproces en dient opgenomen te worden in de technische documentatie
4. Aflevering van de producten in een veilige standaard configuratie



Fook Hwa Tan

Leo van Koppen

Het moment om samen de schouders eronder te zetten en een betrouwbaar digitaal landschap te creëren

5. Producten moeten een conformiteitsprocedure ondergaan en met een CE-markering aantonen te voldoen aan de security eisen
6. Er dient een periode gedefinieerd te worden waarin beveiligingsupdates worden verstrekt
7. Melding van actief misbruikte kwetsbaarheden
8. Importeurs en distributeurs moeten ervoor zorgen dat producten aan de security eisen voldoen als ze op de Europese markt komen

Ik verwacht dat de wet de nodige impact zal hebben voor producenten en leveranciers, met name bij de implementatie in hun ontwikkelprocessen, de meldplicht en de update-garantietijd. Overtredingen kunnen leiden tot boetes van maximaal 15 miljoen euro of 2,5% van de jaarlijkse wereldwijde omzet, afhankelijk van wat hoger is.

Hoewel dit voor sommige producenten een uitdaging zal zijn, biedt de CRA ook een kans om vertrouwen op te bouwen bij klanten en concurrentievoordeel te behalen. Een déjà vu naar de invoering van de GDPR? Ik ben niet alleen geïnteresseerd wat de redacteurs ervan vinden, maar ben eigenlijk nog meer benieuwd wat u als lezer hiervan vindt.

Fook Hwa Tan – Kwaliteit in een nieuw tijdperk: hoe de Cyber Resilience Act hoop biedt voor softwareontwikkeling

De invoering van de Cyber Resilience Act (CRA) is een doorbraak. Het is niet zomaar een stap, maar een mijlpaal op weg naar een digitaal Europa dat veiliger, transparanter en verantwoordelijker is. Voor sommigen misschien klein nieuws, maar voor wie het belang van kwaliteit begrijpt, is dit groot – revolutionair zelfs. Neemt niet weg, dat onze eindbestemming bereiken nog wel wat tijd gaat kosten.

Eindelijk wordt de technologie-industrie tot actie gemaand. De CRA gaat verder dan het opleggen van regels; het vraagt om een fundamentele cultuurverandering. Het dwingt ons de vraag te stellen: wat betekent het om verantwoordelijkheid te nemen voor de digitale producten die ons leven doordringen? Hoe zorgen we ervoor

dat kwaliteit en veiligheid geen losse eindjes zijn, maar vanaf de eerste schets tot de laatste update stevig verankerd zijn in elke fase van ontwikkeling?

En dit is nog maar het begin. De CRA omarmt een holistische aanpak die verder reikt dan softwarekwaliteit alleen. Het pakt de kern van digitale veiligheid aan door:

- Beveiliging te eisen vanaf de ontwerpfase (Security by Design) en veilige standaardconfiguraties te verplichten
- Interfaces en aanvalsoppervlakken te beperken, waardoor cyberaanvallen minder kans krijgen
- Risicobeoordelingen en technische documentatie te verplichten – geen vage beloftes meer, maar harde, toetsbare feiten
- Fabrikanten te dwingen beveiligingsupdates te garanderen en kwetsbaarheden te melden, om vertrouwen en transparantie te bevorderen
- Importeurs en distributeurs medeverantwoordelijk te maken voor naleving, waardoor de hele keten scherp blijft

Ja, de CRA vraagt veel. Het zal bedrijven dwingen om hun processen opnieuw uit te vinden. Het betekent werk. Het betekent verandering. Maar laten we eerlijk zijn: is dat niet precies wat nodig is? Is het niet tijd dat we stoppen met het accepteren van halve oplossingen en ons richten op echte vooruitgang?

De CRA is niet alleen een verplichting, maar ook een kans. Een kans om te excelleren. Een kans om vertrouwen te winnen. Een kans om technologie echt in dienst te stellen van mensen, zonder ze bloot te stellen aan onnodige risico's.

Het is een belofte. Een belofte van een digitale toekomst waarin veiligheid de norm is, waarin kwaliteit geen optie is, maar een fundamenteel recht. Dit is het moment om samen de schouders eronder te zetten en een digitaal landschap te creëren waar we allemaal op kunnen vertrouwen. Wat denkt u? Sluit u aan en help ons de lat hoger te leggen. De tijd van half werk is voorbij – samen bouwen we aan een veilige, rechtvaardige digitale wereld.



Martijn Hoogesteger is Head of Cyber bij S-RM en is bereikbaar via m.hoogesteger@s-rminform.com.

Kwantummechanica uitleggen aan je hond

Ik blijf het maar herhalen: de cybersecuritywereld verandert snel, is gigantisch dynamisch. Bijblijven kan lastig zijn. Verandert Microsoft voor de tiende keer de naam van hun security tools? Jij moet op de hoogte zijn. Is er weer nieuwe wetgeving op Europees, nationaal of sectoraal niveau? Die NIS2 toepasbaarheid heb jij natuurlijk al getoetst. Is er een nieuwe soort phishing die wordt ingezet? De training is al aangepast. Om te zorgen dat we veilig blijven, moeten we in ieder geval bijblijven, en dat kan vaak lastig zijn.

Ik ga de lat nog hoger leggen. Het is naar mijn mening niet genoeg om reactief bezig te zijn. Als we reactief zijn, lopen we altijd achter in de wapenwedloop. Een wapenwedloop die we langzaam beginnen in te halen en ik wil ons weleens voorop zien lopen! Dat betekent onderliggende principes beter begrijpen, en de tijd te hebben om na te denken hoe die principes toepasbaar zijn op de organisaties die we proberen te verdedigen.

Stel, in 2030 is er een statelijke actor die beschikt over een quantumcomputer die klassieke encryptie kan kraken. Prima, denk je. Dat is pas over een lange tijd! Maar stel, die statelijke actor slaat nu al versleutelde data op om het in de toekomst te kraken. Moeten je geheimen over vijf jaar nog steeds geheim zijn? Waar staat al die data? Over welke versleutelde kanalen wordt dat allemaal verstuurd? Is die versleuteling kraakbaar? Hopelijk zijn het onderwerpen waar je al een begin mee hebt kunnen maken, want die actoren zijn inderdaad al versleutelde data aan het verzamelen.

Wat denk je dat er in een NIS3 komt te staan? Wat zouden mogelijke effecten kunnen zijn van geopolitieke spanningen die verder oplopen? Is AI extra relevant de komende jaren of zijn LLMs gewoon een hype? Bestaat het internet over tien jaar nog? En zijn we dan eindelijk over naar IPv6? Zijn er risico's verbonden aan photonic computing? Om goed antwoord te geven moeten we experts worden in IT, kwantummechanica, internationale politiek en het liefst ook wat van risico-afweging snappen, toch? Risico's, die snappen we in ieder geval. En van de toekomst kun je ook een risicoschatting maken! Welke onderwerpen, schat je in, moet je induiken, en welke zijn misschien minder belangrijk? Wat zorgt ervoor dat jij een goede expert op het gebied van informatiebeveiliging blijft, en bij welke onderwerpen loop je het risico dat je achter komt te liggen in de wapenwedloop?

Quantum zal gegarandeerd zo'n onderwerp zijn, dus lees je eens in. Dat kan ook op leuke manieren, zo is 'How to teach Quantum Physics to your dog' (1) een van mijn favoriete boeken. Heb je ook gelijk een vorm van uitleg die de rest van de organisatie wellicht ook snapt.

Ik hoop dat ik de afgelopen jaren een aantal onderwerpen op je radar heb mogen zetten middels deze column. Het is helaas mijn laatste in dit blad. Succes in je risicoreis, en probeer de nieuwe onderwerpen te blijven spotten!

Referentie

(1) How to teach Quantum Physics to your dog – Chad Orzel



CET® Preparation Course

- Doe gedegen kennis op van de vier CET-domeinen: Cloud, Blockchain, IoT en Artificial Intelligence
- Krijg inzicht in de belangrijkste opkomende technologieën
- Met vele praktische oefeningen
- De perfecte voorbereiding voor het internationale CET®-examen



CISM® Preparation Course

- Doe gedegen kennis op van de vier CISM®-domeinen
- Leer hoe je deze domeinen in de praktijk kunt toepassen
- Inclusief praktijkcasussen en een uitgebreid CISM®-proefexamen
- De perfecte voorbereiding voor het internationale CISM®-examen



Scan de QR code(s) voor meer informatie > Of ga naar: www.securityacademy.nl



COLOFON

IB Magazine is het huisorgaan van het Platform voor InformatieBeveiliging (PvIB) en bevat ontwikkelingen en achtergronden over onderwerpen op het gebied van informatiebeveiliging.

HOOFDREDACTEUR

Chris de Vries

REDACTIE

Bianca Brooijmans
Alex Dingemanse
Marcel Feenstra
Maarten Hartsuijker
Fook Hwa Tan
Lilian Knippenberg
Leo van Koppen
Rachel Marbus
Chris de Vries

BLADMANAGEMENT

MOS bv
Caroline Knobbe
Sam Dekkers
E ibmagazine@pvib.nl

ADVERTENTIE-ACQUISITIE

MOS bv
Eric Noordam
E acquisitie@mos-net.nl
T 033 247 34 00

VORMGEVING

Neverseen Art & Design
Dimitri van den Berg

DRUK

Veldhuis Media, Meppel

UITGEVER

Platform voor InformatieBeveiliging (PvIB)
Postbus 1058
3860 BB NIJKERK
T (033) 247 34 92
E secretariaat@pvib.nl
W www.pvib.nl

ABONNEMENTEN

De abonnementsprijs in 2024 bedraagt
€ 118,50 (exclusief btw), prijswijzigingen
voorbehouden.

ABONNEMENTENADMINISTRATIE

Platform voor InformatieBeveiliging (PvIB)
Postbus 1058
3860 BB NIJKERK
E secretariaat@pvib.nl



Tenzij anders vermeld valt de inhoud van dit
tijdschrift onder een Creative Commons
Naamsvermelding-gelijkeDelen 3.0 Nederland
Licentie (CC BY-SA 3.0)
ISSN 1569-1063

Ransomware? NIS2?

SPEED UP YOUR CYBERSECURITY CAREER IN 2025

*Want security
start bij mensen!!*



TSTC

ICT en Security Trainingen

- AIGP** Certified AI Governance Professional
- CNIS2** Certified NIS2 Professional
- BIO** Certified Bio Professional Foundation/Practitioner
- OSCP** OffSec PEN-200
- OSEE** OffSec EXP-401
- CEH** Certified Ethical Hacker v13 AI

TSTC is een gerenommeerd IT opleidingsinstituut en erkend specialist in informatiebeveiliging en cybersecurity trainingen.

GET SKILLED
WWW.TSTC.NL

TECHNICAL SECURITY TRAININGEN

- CEH** Certified Ethical Hacker
- CHFI** Computer Hacking Forensic Investigator
- CPENT** Certified Penetration Testing Professional
- SSCP** Systems Security Certified Professional
- OSCP** Offensive Security Certified Professional

SECURITY MANAGEMENT TRAININGEN

- CISSP** Certified Information Systems Security Professional
- CISM** Certified Information Security Manager
- CISA** Certified Information Systems Auditor
- CRISC** Certified In Risk And Information Systems Control
- CJCSO** Certified Chief Information Security Officer

PRIVACY TRAININGEN

- CIPP/E** Certified Information Privacy Professional / Europe
- CIPM** Certified Information Privacy Manager
- CIPT** Certified Information Privacy Technologist
- CDPO** Certified Data Protection Officer

CLOUD SECURITY TRAININGEN

- CCSP** Certified Cloud Security Professional

ISO TRAININGEN

- ISO 27001** Foundation
- ISO 27001** Lead Implementer
- ISO 27001** Lead Auditor
- ISO 27005** Risk Manager
- ISO 27001** Privacy Management

ONZE TRAININGEN ZIJN KLASSIKAAL OF LIVE ONLINE TE VOLGEN