

Full paper

AI en ARBO; hoe wordt er over gedacht en hoe gaan we vooruit?

Coen van Gulijk¹

Trefwoorden: AI, beleid, forecasting

Samenvatting

Dit artikel rapporteert over een toekomstverkenning over het effect van digitalisering en AI op arbeidsomstandigheden en -veiligheid. Het focust op de visie van stakeholders en hoe zij denken dat AI zijn weerslag zal hebben op arbeidsomstandigheden en -veiligheid. Er zijn semigestructureerde interviews uitgevoerd volgens de zeven-vragen methode. Deze techniek levert een breed palet aan zienswijzen op op het gebied van techniek en beleid. Een belangrijk inzicht is dat voornamelijk risico's AI worden gezien. Van de overheid wordt een veel actievere houding verwacht om in het publieke debat te stappen maar ook om samen met stakeholders te experimenteren, bijvoorbeeld in juridische experimenten. In breder zin is de conclusie dat de discussie over AI in relatie tot veiligheid en gezondheid in het kader van arbobeleid bij alle partijen in Nederland nog in de kinderschoenen staat. En dat komt niet door een gebrek aan vooruitgang in de wetenschap.

Introductie

Nederland is technologische hardloper en stuit daarom snel op kansen en risico's van nieuwe technologieën. Als het gaat om de opmars van kunstmatige intelligentie (AI) levert het een breed spectrum van heel verschillende ARBO-risico's op. Digitalisering en AI leveren aan de ene kant veiligheidswinst op doordat menselijke handelingen kunnen worden vervangen of ondersteund bij repetitieve, gevaarlijke of vieze arbeidstaken, maar tegelijkertijd is er nog veel onduidelijk over de nieuwe gevaren die de technologie oplevert.

Dit artikel rapporteert over een toekomstverkenning naar de effecten van digitalisering en AI voor arbeidsomstandigheden in Nederland. Daarbij beperkt de inventarisatie zich tot de 'smalle' interpretatie van de effecten van digitalisering en AI op arbeidsomstandigheden: het effect van AI op veiligheid in relatie tot de arbeidsomstandigheden. Daarbij is gezocht naar overzicht door verschillende stakeholders om hun zorgen en dromen te delen in een toekomstverkenning over de relatie tussen AI en ARBO. Hierbij werd de 'zeven-vragentechniek' gebruikt. In dit onderzoek werd gewerkt met de onderzoeksvraag:

Abstract

This paper reports a horizon scanning exercise about the effects that digitalization and AI may have on the Health and Safety policy domain. The work focuses on the visions for the future that various stakeholders have and how they think the domain should develop. The research was done using the seven questions method from The Futures Method which is frequently used for policy making in the United Kingdom. The results show that AI and digitalisation are perceived as negative influences on OSH. Government could take a more active role in the public debate and partake in experiments to test policy options (such as legal experimentation). In many ways the work shows that the discussion about AI in the HSE domain is in its infancy in the Netherlands even if scientific progress was made for a more mature discussion.

Hoe denken verschillende stakeholders over de rol van digitalisering en/of AI op arbeidsveiligheid en/of arbeidsomstandigheden en wat vinden zij belangrijk voor de ontwikkeling van het beleidsdomein?

Beknopte achtergrond

Kunstmatige Intelligentie of AI is een verzamelnaam voor verschillende mathematische computertechnieken die, mits op de juiste manier getraind, cognitieve functies van mensen imiteert. Het standaardwerk van Russel en Norvig (2011) geeft aan dat (kunstmatige) intelligentie zich richt op het imiteren van rationele actie. Daarin onderscheiden zij zes relevante capaciteiten: tekstanalyse, kennismodellering, automatisch redeneren, machinaal leren, computervisie en robotica. Het gaat te ver voor dit artikel om de wetenschappelijke vooruitgang op vlakken te behandelen maar het is wel duidelijk dat dit type technologische ontwikkelingen nieuwe (arbeids)risico's oplevert. Voorbeelden daarvan zijn: botsingsrisico's door zelfrijdende goederenrobots (in het Engels: Automated Guided Vehicles of AGV's; bijvoorbeeld in: Cheng et al., 2021), meer fysieke belasting voor medewerkers in de agrarische sector (Benos & Bochtis, 2022), techno-stress in kantooromgevingen (Bonanini et

¹ TNO, University of Huddersfield, Technische Universiteit Delft; email: coen.vangulijk@tno.nl
Correspondentie adres: TNO Healthy Living & Work, Sylviusweg 71, 2333BE Leiden

al., 2020) en slechte werkomstandigheden voor platform-medewerkers (zie o.a. Moore, 2019). Terwijl het onderzoek naar negatieve veiligheidseffecten van technologisering en AI voortzet kijken we in dit onderzoek naar hoe verschillende stakeholders in Nederland de toekomst zien in relatie tot de impact van digitalisering en/of AI op arbeidsveiligheid.

Methode

De werkmethode voor dit onderzoek is gebaseerd op een gestandaardiseerde methode voor beleidsonderzoek: de 'The Futures Toolkit'. Dit is de standaardwerkmethode voor de Engelse overheid voor de ontwikkeling van wetgeving en beleid waarbij toekomstverkenning, draagvlak, en het toetsen van voorgenomen beleid centraal staan. Dit artikel beschrijft één van de methoden uit de toolkit, de zeven-vragen methode. Deze methode verzamelt zienswijzen door middel van semigestructureerde interviews.

Zeven-vragen methode

De zeven-vragen methode richt zich op het bespreekbaar maken van strategische thema's. De methode is geschikt om controversen te identificeren en kan in zowel groepen als met individuen worden uitgevoerd. In dit onderzoek is ervoor gekozen om individuele interviews te doen om te voorkomen dat experts elkaar beïnvloeden en om een zo breed mogelijk spectrum van onderwerpen te verzamelen. Er is geen maximum aan het aantal mensen dat op deze manier kan worden geïnterviewd. Dit onderzoek was afgebakend door 12 experts te interviewen; hiermee volgen we de richtlijnen van Hennink & Kaiser (2021) die suggereren dat bij dit aantal gemiddeld genomen een verzadiging optreedt in de hoeveelheid nieuwe informatie die verkregen wordt uit de interviewdata (en dat daarmee de efficiëntie voor informatie-extractie voor additionele interviews laag wordt).

Selectiecriteria voor deelnemers waren de volgende: i) bewezen ervaring met het onderwerp welke werd geverifieerd door publicaties. Dit leverde in de meeste gevallen ook de contactgegevens op ii) vertegenwoordiging van een unieke belangen- of kennegroep. Hierbij was het vooral belangrijk dat het spectrum aan instituten zo breed mogelijk te maken. iii) Nederlandsprekend. Dit om de focus op Nederlandse situatie te leggen. Het laatste criterium is een relatief grofstoffelijke manier om de focus op Nederlandse zienswijzen te leggen. De 12 experts vanuit verschillende organisaties: vakbonden, bedrijven, (veiligheids-)onderzoeksinstituten, universiteiten en ondernemersorganisaties. Twee interviews waren met Nederlandstalige individuen die werken bij internationale instituten. Merk niet elke geïnterviewde die Nederlands spreekt in Nederland woont en ook niet automatisch van Nederlandse afkomst is. De interviews zijn afgenomen via beeldbellen.

De vragen waren altijd dezelfde, namelijk:

1. Als u met iemand zou kunnen praten die 10 jaar in de toekomst leeft; wat zou u die persoon vragen over wat digitalisering en AI hebben betekend voor veiligheid en gezondheid (ARBO)?

2. Stel nu dat de persoon een succesverhaal deelt, hoe ziet dat succes er volgens u dan uit?
3. Wat zouden de bedreigingen zijn voor het bereiken van uw visie op succes?
4. Wat is de beste manier om huidig ARBO-beleid te veranderen om uw visie te bereiken?
5. Nu kijken we terug naar de afgelopen 10 jaar; van welke specifieke successen en/of missers kunnen we leren?
6. En wat moet er NU worden gedaan om uw visie van succes te bereiken?
7. En tenslotte, als u het voor het zeggen had en alles mogelijk zou zijn, welke eerste stap zou u dan zetten?

Van de interviews zijn verslagen geschreven die aan de geïnterviewden werden voorgelegd zodat zij de vrijheid hadden om het verslag naar eigen inzien aan te passen. Deze verslagen vormen de brondata voor het onderzoek.

Bewerking van de data: amalgamatie en aggregatie

De brondata zijn door amalgamering en aggregatie teruggebracht naar zeven uitspraken. Daarmee volgen we de suggestie die de 'The Futures Toolkit' geeft om de antwoorden te formuleren alsof de zeven vragen worden beantwoord door verschillende personae die je kunt associëren met de zeven vragen: een helderziende beantwoordt vraag 1, een optimist vraag 2, een pessimist vraag 3, een visionair vraag 4, een historica vraag 5, vraag 6 een activiste, en vraag 7 wordt beantwoord door een strateeg. Dit levert tot op zekere hoogte een categorisering op van de antwoorden.

Er is een aantal stappen nodig om van interviewverslagen tot uitspraken van personae te komen. De volgende stappen zijn gemaakt:

1. Antwoorden van alle respondenten op de vragen amalgameren per vraag in een tabel.
2. Aggregatie van uitspraken per vraag, dubbelures samenvoegen, waar nodig onderwerpen onderverdelen in lijsten. Dit zijn de resultaten als lijst van uitspraken.
3. Samenvatten van bevindingen in één antwoord per vraag: de kern-resultaten. Dit is de stap waar de onderzoekers de resultaten het sterkst beïnvloeden doordat zij interpretaties moeten doen.

Resultaten

De resultaten zijn weergegeven in de vorm van een uitspraak vanuit het perspectief van één persona.

De **historica** vertelt je dat de adoptie van digitalisering en AI in de afgelopen 10 jaar veel langzamer is verlopen dan 10 jaar geleden gedacht werd. Dat kwam niet omdat wetten ontwikkeling in de weg zouden staan, maar omdat zelfregulering in dat opzicht mislukt is: digitalisering heeft ons wel heel druk gemaakt, maar het werk is niet beter of leuker geworden, met algoritmisch management als belangrijkste uitwas. De autoriteiten kunnen helpen door een duidelijker narratief op te stellen om bedrijven uit te leggen hoe het wettelijk kader te gebruiken bij innovatie (maar de historica kan geen voorbeeld geven), en actief

deel te nemen aan die innovatie. Als er dan innovaties worden gedaan, ontbreekt het helaas aan reflectie door de afwezigheid van (langdurige) monitoring van arbeidsomstandigheden, bewijs van (on)veiligheid, ongevallenstatistieken en benchmarking wat slechte veiligheidsproducten en -interventies oplevert.

De **helderziende** vertelt dat in de komende 10 jaar waarschijnlijk veel verandert: sommige arbeidsprocessen gaan onherkenbaar veranderen, beroepsgroepen komen en gaan, competenties veranderen, en er zullen verschuivingen zijn in waar werken het leukste is. Ze vertelt dat number-crunching en AI het werk van veiligheidskundigen gemakkelijker gaan maken en dat er minder ongevallen zullen zijn. Fysieke klachten, stress en burnouts verminderen tegelijkertijd. Ze kan je ook uitleggen dat (maar nog niet hoe) de balans tussen zorgplicht van de werkgever en rechten van de medewerker is verschoven onder invloed van economische productieoptimalisatie, arbeidsflexibiliteit en (vooral psychosociale) gezondheid en privacy.

De **optimist** vertelt je dat werken gezonder én leuker wordt. Mensen blijven geestelijk scherp en fysiek sterker. Repetitief, vies, en zwaar werk wordt overgenomen door robots. Kwetsbare arbeidsgroepen kunnen goed mee in de moderne tijd en er wordt meer uitdagend onderhoud gepleegd. Door veiligheidsinnovaties is werk inherent veiliger. AI verbetert RI&E, arbeidscondities, inspraakprocedures, arbeidshygiënische strategie, en vermindert regeldruk. Baanzekerheid is vervangen door een breed spectrum van arbeidsaanbieders waar werknemers uit kunnen kiezen.

De **pessimist** vertelt je dat AI-innovatie is ingegeven door technologie-push en durfkapitalisme (venture capitalism). Door gebrek aan toetsing aan maatschappelijke behoefte blijven we steken in het efficiënter maken van kort cyclisch repetitief werk in slechte condities zoals scannen of klikken. Door gebrek aan competentie in het veiligheidsdomein kunnen stress en hoge werkdruk nauwelijks beheerst worden en vooral ouderen vallen uit. Daarmee blijft een verandercultuur uit en blijven medezeggenschap en werknemersparticipatie buiten zicht; het gebrek aan privacy helpt daar niet bij.

De **visionair** stelt dat bestaande wetten en systeemverplichtingen geen beperkingen leveren voor innovatie en dat er behoefte is aan een overheid die AI-innovatie aanjaagt. Dat kan de overheid doen door fondsen voor specifieke thema's vrij te maken, door zélf geavanceerde technieken toe te passen, door informatiebronnen of standaarden daarvoor beschikbaar te stellen, en door betrokkenheid met lopende onderzoeken. Aandachtspunten in een op te stellen roadmap zouden moeten zijn: mens-machine gedrag, frequente (maar risicovolle) werkzaamheden, bronaanpak, bewustwording & RI&E, en juridische experimentatie.

De **activiste** vindt dat de overheid de regie op ARBO moet

terugpakken door de inhoudelijke ambtenaar terug te halen en onderzoeksgeld vrij te maken voor onafhankelijk onderzoek naar de schadelijke gevolgen van AI op arbeid. Tegelijkertijd wil ze dat de overheid strengere veiligheidseisen opstelt en die strenger gaat controleren. Ook wil ze dat bedrijven de cultuur binnen het bedrijf klaar maken voor de toekomst (maar kan nog niet zeggen hoe dat moet). Daarnaast wil zij fundamenteel en veldonderzoek om de gevolgen van AI op ARBO scherper te maken. Hiervoor wil zij dat we helemaal uit de box gaan denken: als we alles helemaal opnieuw zouden inrichten met AI, wat krijgen we dan?

De **strateeg** zegt drie dingen. Ten éérste: we moeten af van de bijna verborgen insluiping van AI in het arbeidsproces. Creëer zichtbaarheid van de technologische ontwikkelingen en maak een overzicht over de impact. Blijf de ontwikkelingen over een langere periode volgen. Verzamel de succesverhalen en verander daarmee het ARBO-systeem waar nodig. De tweede: ga experimenteren! Ga aan de slag met nieuwe ideeën, van verkennend onderzoek tot praktische toepassing van systemen die in andere werelddelen al (mogen) worden ingezet en alles wat daartussen zit. Ten derde: het lerend vermogen van bedrijven en de overheid moet vergroot worden door proactieve betrokkenheid met AI, ARBO en de relevante stakeholders.

Discussie

Diepgang en bruikbaarheid van de resultaten.

De resultaten zijn geformuleerd als uitspraken. De uitspraken zijn samenvattingen waarbij getracht is de stem van de deelnemers zo sterk mogelijk te laten doorklinken. Dat levert kort maar krachtige statements op die verder gaan dan alleen het opsommen van onderwerpen of problemen; het geeft ook inzicht in wat deelnemers vinden van de problematiek en soms ook hoe die op te lossen (maar dit laatste helaas niet in groot detail). Het voordeel is dat daarmee 'diepte' wordt gecreëerd, tegelijkertijd introduceert het mogelijkheden om het oneens te zijn met de zienswijzen van de deelnemers en is er tegenspraak in de uitspraken mogelijk. Bijvoorbeeld, de visionair stelt dat huidige wetten en systeemverplichtingen geen beperking opleveren voor innovatie. De pessimist vindt dat het gebrek aan privacy voortgang in de weg staat en daarmee ontstaat een tegenstelling. In dit werk zoeken we niet uit wie van beiden gelijk heeft (als de vergelijking überhaupt mogelijk is) maar we nemen het als gegeven mee dat er verschillende zienswijzen zijn en dat die aanleiding zijn voor verdere verdieping in vervolgonderzoek.

De uitspraken die gemaakt zijn, zijn niet op feitelijke juistheid getoetst. Wellicht de geïnterviewden die bij hebben gedragen aan het statement van de helderziende een goed inzicht in hoe 'number-crunching' het werk van veiligheidskundigen gemakkelijker gaan maken; maar die is dan voor de meeste geïnterviewden onvoldoende zichtbaar en dus alléén kort genoemd door de helderziende.

Zeven-vragen methode

Met de zeven-vragen methode kan een zekere mate van diepgang bereikt worden. Waar toekomstverkenningen vaak opsommingen zijn van onderwerpen kunnen we nu ook aangeven wat mensen ervan vinden en wat (en of) ze er aan zouden willen doen. Dat is heel duidelijk bij het perspectief van de pessimist: die vermoedt dat AI leidt tot hoge werkdruk en -stress. Door het samenbrengen van een constatering en een mening wordt meer diepgang bereikt. Dit maakt het gemakkelijker om uitgangspunten voor verandermanagement of (overheids-)beleid op te stellen doordat niet allen problemen maar ook oplossingen bespreekbaar worden, alsook de reden waarom mensen denken dat die oplossingen relevant zijn.

De zeven-vragen methode is een efficiënte manier om relevante informatie uit de markt te halen; de interviews zijn relatief kort (meestal is 30 minuten genoeg), de vragen zijn van te voren duidelijk en de verwerking naar rapportage is relatief eenvoudig. Daarnaast is de methode beproefd en (in het Verenigd Koninkrijk) vaak toegepast.

Antwoorden

De onderzoeksvraag van dit werk was de volgende:

Hoe denken verschillende stakeholders over de rol van digitalisering en/of AI op arbeidsveiligheid en/of ARBO en wat vinden zij belangrijk voor de ontwikkeling van het beleidsdomein?

Deze vragen worden beantwoord vanuit de resultaten in paragraaf 3. De antwoorden worden gegeven als een korte lijst van bevindingen die zijn onderverdeeld in 5 categorieën: brede sociale vraagstukken, de rol van de overheid, voorstellen voor systeemverandering, kansen en risico's van de techniek, problemen van de techniek en dwarse ideeën. De laatste categorie bevat uitspraken die als 'out-of-the-box' kunnen worden gezien. De antwoorden zijn:

- De stakeholders vinden dat er een bredere sociale discussie nodig is over de condities waaronder digitalisering en AI in het arbeidsproces worden ingezet. Verschillende thema's zijn daarbij belangrijk: kwaliteit en kwantiteit van werk, medezeggenschap bij innovatie en installatie, en persoonlijke versus collectieve belangen.
- Ze vinden dat de overheid een rol heeft in de regie van AI-innovatie op de werkplek; ze kan innovatiekaders (beter) omschrijven; ze kan innovatie aanjagen; en ze kan zelf deelnemen aan innovatietrajecten.
- De systeemaanpak voor het arbeidsomstandighedenbeleid heeft aandachtspunten in relatie tot AI en digitalisering, met name in relatie tot aandacht voor de thema's welzijn en psychosociale belasting, competenties, cultuur, en samenwerking in de keten. Daarnaast ook aandacht voor het opstellen van een technologische roadmap, en het lerend vermogen van overheid en bedrijven vergroten.

- AI heeft de potentie om werk veiliger en leuker te maken en het werk van veiligheidkundigen gemakkelijker te maken, maar deze beloften zijn nog niet waargemaakt. Tot nu toe heeft AI vooral schade gedaan aan werknemers wat het vertrouwen in AI niet ten goede komt (zie ook Benos et al., 2021; Moore, 2019).
- Tenslotte gaan er veel stemmen op voor de exploratie van radicale oplossingen om AI veilig in het arbeidsveiligheidsdomein te brengen. Het loont de moeite om een aantal paden te onderzoeken.

Onzichtbare innovatie

Dit artikel laat zien dat er vooralsnog vooral aan de negatieve effecten van AI wordt gedacht als het gaat om arbeidsomstandigheden en arbeidsveiligheid en dat positieve effecten niet merkbaar zijn. Dat is jammer, want er wordt op verschillende vlakken aan veiligheidsinnovaties gewerkt met AI. In deze paragraaf een paar woorden daarover.

Tixier was in 2016 één van de eerste auteurs die rapporteerde over succesvol gebruik van tekstanalyse van bouwongeval-rapportages in de Verenigde Staten (Tixier et al., 2016). De auteur van dit werk koppelde in 2019 tekstanalyses aan bestaande BowTies voor het Engelse spoor (Hughes et al., 2019), en onderzocht de toekomst van veiligheid met industriepartijen (IEC, 2020) en academici (Swuste et al. 2020). En het onderwerp blijft onverminderd relevant: zie bijvoorbeeld Mutlu et al. (2023) die breder kijkt naar alle arbeidsongevallen in Turkije. Inmiddels wordt ChatGPT breed besproken in het veiligheidsdomein.

Vrij recent worden vele verschillende AI-technieken ingezet voor een betere werkomgeving. Johnson (2022) gebruikt lerende algoritmes om faalmechanismen in mens-machine interacties met slimme machines op te sporen en geeft ook oplossingen om dat te verbeteren. Guzman et al. (2016) gebruiken kunstmatige intelligentie kennismodellen om risicoanalyses te verbeteren. En Liu et al. (2022) gebruiken meerdere AI-technieken om tot een soort automatisch redeneren te komen voor risico's van seinen op het spoor.

Ook tekent zich nu onderzoek af dat in verband kan worden gebracht met nieuwe Europese wetgeving op dit vlak. Anastasi et al. (2021) proberen met machinaal leren patronen te herkennen in arbeidsongevallen om iets te kunnen zeggen over de essentiële veiligheids- en gezondheidseisen zoals die in de Europese machineveiligheidswetgeving voorkomen (EC, 2021a). Braband & Schäbe (2020) en Rudolph et al. (2018) buigen zich over de vraag hoe je een veiligheidsverificatie kunt maken van kunstmatige intelligentie zelf; dit in verband met de in ontwikkeling zijnde AI-act (EC, 2021b). Met name deze laatste categorie zoekt antwoorden die voor beleidsmakers in de overheid en in bedrijven van belang zijn.

Dat de geïnterviewden weinig positieve effecten zien heeft dus niet zozeer te maken met een gebrek aan wetenschappelijke ontwikkelingen en innovaties maar meer dat die

ontwikkelingen de markt onvoldoende bereiken. Hier ligt dus nog een uitdaging voor onderzoekers.

Conclusie

Dit artikel rapporteert over een toekomstverkenning over het effect van digitalisering en AI op arbeidsomstandigheden en -veiligheid en focust op de visie van stakeholders en hoe zij denken dat het beleidsdomein zich zou (moeten) ontwikkelen. Voor dit werk is de Engelse 'The Futures Toolkit' gebruikt en daaruit de zeven-vragen methode. Deze techniek levert een breed palet aan zienswijzen op het gebied van techniek en beleid. De resultaten wijzen erop dat de discussie over AI in relatie tot veiligheid en gezondheid in het kader van arbobeleid nog in de kinderschoenen staat. Helaas hebben AI innovaties nog onvoldoende impact op de dagelijkse praktijk gehad. Om dat te verbeteren moet in het ARBO domein meer geëxperimenteerd worden met AI technieken, is een breder(e) sociale discussie nodig en zou het goed zijn als de overheid actief participeert in onderzoek. Dat levert een basis om een beleidsvisie te ontwikkelen in samenspraak tussen voor overheden, bedrijven en medewerkers.

Verantwoording

Dit onderzoek is mogelijk gemaakt door steun van het Ministerie van Sociale Zaken en Werkgelegenheid (SZW) aan het TNO onderzoeksprogramma MAPA Arbeidsveiligheid 2022.

Literatuur

- Anastasi S, Madonna M, & Monica L. (2021) Implications of embedded artificial intelligence - machine learning on safety of machinery. *Procedia Computer Science*, 180: 338-343.
- Benos L, Bochtis D. (2021) An analysis of safety and health issues in agriculture towards work automation. In: Bochtis DD, Pearson S, Lampridi M, Marinoudi V, Pardalos PM. (eds) *Information and Communication Technologies for Agriculture - Theme IV: Actions*. Springer Optimization and Its Applications, vol 185. Springer.
- Bondanini G, Giorgi G, Ariza-Montes A, Vega-Muñoz A, Andreucci-Annunziata P. (2020) Technostress dark side of technology in the workplace: a scientometric analysis. *Int. J. Environ. Res. Public Health*, 17: 8013.
- Braband J, Schäbe H. (2020) On safety assessment of artificial intelligence. *Dependability*, 20 (4): 25-34.
- Cheng Z, Likai J, Min C, Zingyan Z. (2021) Mechanical safety risk analysis of Smart Factory, *Journal of Physics: Conference Series*, Volume 1884, 2021 International Conference on Intelligent Manufacturing and Industrial Automation (CIMIA 2021) 26-28 March 2021, Guilin, China.
- European Commission (EC). (2021a) Proposal for a regulation of the European Parliament and of the council on machinery products. COM(2021) 202 final, 2021/0105(COD). Beschikbaar via: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A52021PC0202>.
- European Commission (EC). (2021b) Regulation of the European Parliament and of the council laying down harmonized rules on artificial intelligence (artificial intelligence Act) and amending certain union legislative acts. Beschikbaar via: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=celex%3A52021PC0206>.
- Ishida GS, Choi E, Aoyama A. (2016) Artificial intelligence improving safety and risk analysis: A comparative analysis for critical infrastructure. 2016 IEEE International Conference on Industrial Engineering and Engineering Management (IEEM), Bali, Indonesia, pp. 471-475.
- Hennink M, Kaiser BN. (2021) Sample sizes for saturation in qualitative research: A systematic review of empirical tests. *Social Science & Medicine*, 292: 114523.
- Hughes P, Shipp D, Figueres M, Van Gulijk C. (2019) From free-text to structured safety management: Introduction of a semi-automated classification method of railway hazard reports to elements on a bow-tie diagram. *Safety Science*, 110 (Part B): 11-19.
- IEC. (2020) IEC White Paper Safety in the future. Beschikbaar via: <https://www.iec.ch/basecamp/safety-future>.
- Johnson B. (2022) Metacognition for artificial intelligence system safety – An approach to safe and desired behavior. *Safety Science*, 151:105743.
- Liu S, Yang S, Chu S, Wang C, Liu R. (2022) Application of Ensemble Learning and Expert Decision in Fuzzy Risk Assessment of Railway Signaling Safety. IEEE 25th International Conference on Intelligent Transportation Systems (ITSC), Macau, China, pp. 3691-3697.
- Moore PV. (2019) OSH and the Future of Work: Benefits and Risks of Artificial Intelligence Tools in Workplaces. In: Duffy, V. (eds) *Digital Human Modeling and Applications in Health, Safety, Ergonomics and Risk Management. Human Body and Motion*. HCII 2019. Lecture Notes in Computer Science(), vol 11581. Springer.
- Mutlu NG, Altuntas S, Dereli T. (2023) The evaluation of occupational accident with sequential pattern mining. *Safety Science*, 166: 106212.
- Russel S. Norvig P. (2016) *Artificial intelligence, a modern approach* 3rd ed. Pearson Education Limited, Essex, England.
- Swuste P, Groeneweg J, Van Gulijk C, Zwaard W, Lemkowitz S, Oostendorp Y. (2020) The future of Safety Science. *Safety Science*, 125: 104593.
- Tixier AJP, Hallowell MR, Rajagopalan B, Bowman D. (2016) Application of machine learning to construction injury prediction. *Automation in Construction*, 69: 102-114.
- UK Government office for Science. (2017) The futures toolkit. Beschikbaar via: <https://assets.publishing.service.gov.uk/media/5a821fdee5274a2e8ab579ef/futures-toolkit-edition-1.pdf>.