

AI IN BEELD



Universiteit
Leiden
Geesteswetenschappen

Mogelijkheden, risico's en richtlijnen van
kunstmatige intelligentie voor de beeldende journalistiek

Prof.dr. J.C. de Jong, Dr. A. Vandendaele,
M. van der Woude (MA) & S. Arends (MA)

Leerstoelgroep Journalistiek & Nieuwe Media

Leiden University Centre for Linguistics (LUCL)

© 2025 Universiteit Leiden

Niets uit deze uitgave mag worden vermenigvuldigd en/of openbaar gemaakt door middel van druk, fotokopie of op welke wijze dan ook, zonder voorafgaande schriftelijk toestemming van de auteursrechthebbende

Coverbeeld: Maartje van der Woude, gegenereerd met Midjourney.
Gebruikte prompt: wolf silhouette in Dutch neighborhood street
with cars by night

AI in Beeld

Mogelijkheden, risico's en richtlijnen van
kunstmatige intelligentie voor de beeldende journalistiek

Prof.dr. Jaap de Jong, Dr. Astrid Vandendaele,
Maartje van der Woude (MA) & Stef Arends (MA)

Inhoudsopgave

Voorwoord.....	1
Samenvatting.....	3
1 Inleiding.....	9
1.1 Aanleiding.....	10
1.2 Onderzoeksvragen.....	11
1.3 Methode.....	11
1.4 Uitvoering.....	12
1.5 Opbouw van dit onderzoeksrapport.....	12
2 Literatuurstudie.....	13
2.1 Beeld in journalistiek.....	14
2.2 Objectiviteitsclaim nieuwsfoto's.....	14
2.3 Mogelijkheden van beeld-AI in de journalistiek.....	14
2.3.1 Tijd- en kostenbesparing.....	15
2.3.2 Creativiteit in brainstormfase.....	15
2.3.3 Productie van alternatieve stockfoto's, illustraties, infographics en datavisualisaties.....	15
2.3.4 Visualisatie van niet-fotografeerbare zaken.....	16
2.4 Risico's van beeld-AI in de journalistiek.....	16
2.4.1 Dalend vertrouwen in de geloofwaardigheid van nieuws.....	16
2.4.2 Copyrightproblemen.....	17
2.4.3 AI kan vooroordelen versterken.....	17
2.4.4 Vrees voor baanverlies.....	18
2.4.5 Beperkte AI-expertise en kennissilo's.....	18
2.4.6 Klimaatimpact.....	19
2.5 Journalistieke normen.....	19
2.6 Regelgeving.....	19
2.7 Transparantie.....	21
2.8 AI-tools in de journalistiek.....	22
2.8.1 Technologie in een stroomversnelling.....	22
2.8.2 Beeldgeneratieprogramma's.....	24
2.8.3 Beeldverificatieprogramma's.....	28
2.8.4 Beeldherkenningsprogramma's.....	29
2.8.5 Vier ethische vragen over AI-tools.....	30
3 Methode.....	33
3.1 Deskresearch.....	34
3.2 Diepte-interviews.....	34
3.3 Experiment.....	37
4 Resultaten.....	39
4.1 Mogelijkheden van journalistieke beeld-AI.....	40
4.1.1 'Geslaagde' voorbeelden.....	40
4.1.2 Routinematige efficiëntie.....	44
4.1.3 Mogelijkheden voor illustratie.....	45
4.1.4 Kansen voor 'onmogelijke' beelden.....	46
4.1.5 Kansen voor satire.....	46
4.2 Beeldbewerking versus beeldmanipulatie.....	47
4.3 Ruimte voor experiment op regionale redacties.....	48
4.4 Risico's voor beeld-AI.....	50
4.4.1 Ethische bezwaren.....	50
4.4.2 Bezwaren van 'kennissilo's'.....	53
4.4.3 Verlies van gerespecteerd creatief 'ambacht' van fotograaf.....	54
4.4.4 Baanverlies.....	55
4.5 Richtlijnen.....	55
4.5.1 Weinig richtlijnen.....	55

4.5.2 Richtlijnen zijn noodzakelijk.....	56
4.5.3 Richtlijnen op meerdere niveaus.....	57
4.5.4 Zorgen over externe regelgeving die journalistieke vrijheid en creativiteit beperkt.....	57
4.5.5 Bestaande richtlijnen.....	59
4.5.6 Concrete richtlijnen van respondenten.....	60
4.5.7 De handhaving van richtlijnen.....	60
4.6 Noodzakelijke AI-expertise.....	61
4.6.1 Educatie.....	61
4.6.2 AI-experts.....	62
4.7 Transparantie.....	63
4.7.1 Transparantie essentieel.....	63
4.7.2 Aarzeling over 'te' transparant zijn.....	65
4.8 De waarheid van beelden in de journalistiek.....	66
5 Conclusies en aanbevelingen.....	69
5.1 Conclusies.....	70
5.1.1 AI-beeldtoepassingen worden nog slechts beperkt en voorzichtig ingezet.....	70
5.1.2 AI wordt vooral gebruikt voor routinematige efficiëntie en beeldbewerking.....	70
5.1.3 AI heeft potentie als tool voor illustratieve en abstracte journalistieke visueel toepassingen.....	70
5.1.4 AI maakt het mogelijk om 'onmogelijke' beelden te genereren.....	71
5.1.5 Angst voor baanverlies onder beeldprofessionals.....	71
5.1.6 Journalistieke richtlijnen voor AI-beeldgebruik zijn onvoldoende ontwikkeld.....	71
5.1.7 Er is een gebrek aan AI-expertise en kennisdeling binnen redacties.....	71
5.1.8 Transparantie over AI-gebruik is noodzakelijk maar uitdagend.....	72
5.1.9 AI roept fundamentele vragen op over waarheid en de rol van visuele journalistiek.....	72
5.1.10 AI reproduceert vooroordelen.....	73
5.2 Aanbevelingen.....	73
5.2.1 Ontwikkel en implementeer AI-richtlijnen.....	73
5.2.2 Bied transparantie op maat.....	74
5.2.3 Wees kritisch op door AI gereproduceerde vooroordelen.....	74
5.2.4 Versterk de controle op beeldmanipulatie.....	74
5.2.5 Versterk AI-expertise op de redactie.....	75
5.2.6 Train en informeer redacties over AI-gebruik.....	75
5.2.7 Experimenteer, maar met duidelijke grenzen.....	76
5.2.8 Betrek het publiek in de AI-discussie.....	77
5.2.9 Aanzet concrete richtlijnen.....	77
5.2.10 Vervolgonderzoek.....	79
Referenties.....	80
Bijlage 1 – Overzicht interviews experts.....	85
Bijlage 2 – Overzicht interviews praktijkmensen.....	86
Bijlage 3 – Overzicht AI-tools en hun functionaliteiten.....	87

Voorwoord

AI in Beeld? AI biedt zowel kansen als risico's voor journalistieke beeldproductie. Het vergt nieuwe vaardigheden, transparantie en ethische reflectie van beeldprofessionals. Het succes van AI in de journalistiek zal afhangen van duidelijke richtlijnen, technologische controlemechanismen en een open dialoog met het publiek. Dit onderzoek brengt deze uitdagingen in kaart en reikt oplossingsrichtingen aan.

Het initiatief voor dit project, AI in Beeld, lag bij de Leidse leerstoel Journalistiek en Nieuwe Media. Het Stimuleringsfonds voor de Journalistiek was bereid een groot deel van de kosten te dragen vanuit de Regeling Journalistiek Onderzoek 2023-2024.

Wij bedanken alle negenenvijftig geïnterviewde media-experts, ombudspersonen, foto-journalisten, beeldredacteuren en grafisch vormgevers voor hun tijd, hun betrokkenheid bij het onderwerp en hun ideeën. Ook veel dank aan Tijmen Koppelaar, junior onderzoeker Universiteit Leiden, voor zijn grondige analyses van de interviews (in Atlas.TI), en Marc Prüst (Marc Prüst Consultancy) voor zijn kritische oog en waardevolle suggesties.

Onmisbaar was de samenwerking met de mediapartners in dit project. Algemeen Dagblad, hoofdredacteur Rennie Rijpma en haar chef-beeld Derk Westra, KRO-NCRV's Pointer, namens wie datajournalist Jerry Vermanen een actieve inbreng heeft gehad in dit project en ten slotte de Stichting Regionale Publieke Omroepen met Sharla Hagnaars, senior datamedewerker en lid van het Transformatieteam van de RPO.

Leiden, 15 maart 2025

Jaap de Jong

Astrid Vandendaele

Maartje van der Woude

Stef Arends

Samenvatting

Samenvatting

Met AI-beeldapps is het mogelijk om zonder camera fotorealistische beelden te creëren. Het is voor journalisten problematisch om onder tijdsdruk vast te stellen wat het echtheidsgehalte van dit beeldmateriaal is. Daarnaast roept de toegenomen inzet van AI vragen op over auteursrechten: welke bestaande foto's zijn gebruikt bij de AI-generatie, en hebben de makers daar toestemming voor gegeven of compensatie voor ontvangen? En – misschien erger nog – risico's van het publieke vertrouwen in de journalistiek. In hoeverre blijft journalistiek beeld betrouwbaar als we weten dat het zonder camera kan worden gegenereerd?

Tegelijkertijd zijn er kansen voor de journalistiek: met AI kunnen snel en efficiënt taken worden verricht in de productie en bewerking van beelden. Meer expertise is gewenst bij fotojournalisten en grafisch vormgevers om goed om te gaan met de nieuwe AI-mogelijkheden en bij beeldredacteurs om gemanipuleerd beeld te beoordelen.

Mediaorganisaties moeten dus duidelijke richtlijnen ontwikkelen voor zowel aangeleverd als eigen beeldmateriaal (foto's, illustratie en video). De uitdagingen voor journalistieke media zijn evident en leiden tot de onderzoeksvraag:

Hoe kan de journalistiek op een verantwoorde manier omgaan met de toenemende mogelijkheden en risico's van beeldbewerking en -productie met AI?

In *deskresearch* werd een literatuuronderzoek verricht naar (beeld)richtlijnen in de journalistiek en vertrouwen in de media. Ook is een inventarisatie opgesteld van veelgebruikte visuele AI-applicaties. In dit kwalitatieve onderzoek zijn semigestructureerde diepte-interviews afgenomen met in de eerste fase een diversiteit aan experts. Vervolgens met journalistieke beeldmakers, onder wie fotojournalisten, beeldredacteurs, grafisch vormgevers. Ook zijn ombudspersonen geïnterviewd. In totaal zijn er 19 experts en 40 praktijkmensen als bron gebruikt.

Het initiatief voor dit onderzoek is genomen door de leerstoelgroep Journalistiek en Nieuwe Media van Universiteit Leiden en uitgevoerd in samenwerking met mediapartners Algemeen Dagblad, KRO-NCRV's Pointer en Stichting Regionale Publieke Omroep.

Conclusies

De belangrijkste conclusies van dit onderzoek hebben betrekking op de mogelijkheden, risico's en richtlijnen van beeld-AI in de journalistiek.

Mogelijkheden

Allereerst is het opvallend dat AI-beeldtoepassingen slechts beperkt en voorzichtig worden ingezet in het journalistieke werkveld.

De mogelijkheden van beeld-AI hebben in de praktijk vooral betrekking op routinematige efficiëntie en beeldbewerking. AI blijkt een waardevol hulpmiddel bij het uitvoeren van repetitieve taken, zoals het snel genereren van schetsen, eenvoudige visuals of concepten. Ook wordt AI in de praktijk vaak gebruikt voor beeldbewerkingstaken zoals het aanpassen van kleuren, contrast en scherpte.

Samenvatting

Daarnaast heeft AI potentie als tool voor illustratieve en abstracte journalistieke visuele toepassingen, met name in situaties waarin geen geschikt beeldmateriaal beschikbaar is. AI kan dienen als een kostenbesparend alternatief voor stockfoto's of handgemaakte illustraties, of wanneer complexe fenomenen verbeeld moeten worden. Verder wordt AI ingezet voor het genereren van visuals bij podcasts, explainer-video's. AI maakt het verder mogelijk om 'onmogelijke' beelden te genereren, in situaties waarin traditionele beeldproductie niet haalbaar is vanwege ethische, juridische of veiligheidsbeperkingen. Tot slot wordt beeld-AI ingezet in media waarin satire en parodie een rol spelen.

Opvallend is dat regionale journalistieke media meer ruimte lijken te hebben voor experimenten met beeld-AI dan hun nationale tegenhangers. Het bekendste Nederlandse voorbeeld is een experiment van Omroep Brabant, waarin de AI-kloon van een menselijke nieuwslezer een paar maanden het nieuws heeft gepresenteerd.

Risico's

Beeldproductie met AI roept vragen op over auteursrechten. Veel AI-modellen worden getraind op grote datasets, waarin ook auteursrechtelijk beschermde werken kunnen zitten. Het is vaak onduidelijk of beeldmakers toestemming hebben gegeven of vergoed worden voor hun werk. Dit leidt tot juridische en ethische discussies over eerlijke compensatie van beeldmakers. Een ander ethisch risico is het misbruik van AI voor het creëren van deepfakes, waarmee misleidende of schadelijke beelden en audio eenvoudig kunnen worden geproduceerd. De journalistiek heeft te maken met de circulatie van nepbeelden en moet (digitale) manieren vinden om die te onderscheiden van echte. Wanneer media zelf AI-beeld publiceren, kan het publiek misleid worden door de suggestie van integer, fotorealistisch beeld. Zelfs als er onderschriften en disclaimers toegevoegd zijn, bestaat het risico dat publiek dat niet opmerkt.

Veel respondenten tonen grote bezorgdheid over het bredere maatschappelijke effect van AI-gegenereerde, fotorealistische beelden op het vertrouwen in de journalistiek. Omdat AI-modellen zich trainen op materiaal waarin onvermijdelijk vooroordelen voorkomen, produceren ze eindproducten met daarin stereotypen over bijvoorbeeld geslacht en huidskleur.

Baankansen van beeldprofessionals nemen door de opkomst van AI af, verwacht een aantal respondenten. Dat komt volgens hen door AI-toepassingen die routinematige efficiëntie vergroten.

Richtlijnen

Journalistieke richtlijnen voor AI-beeldgebruik zijn nog onvoldoende ontwikkeld. Experts pleiten voor redactierichtlijnen én sectorbrede regulering door beroepsverenigingen, zoals de NVJ en de Raad voor de Journalistiek. Wat betreft richtlijnen benadrukken beeldredacteuren en fotografen vooral het belang van waarborging van journalistieke vrijheid.

Samenvatting

Aanbevelingen

Op basis van deze resultaten komen de onderzoekers tot de volgende aanbevelingen voor de journalistiek.

Ontwikkel en implementeer AI-richtlijnen. Laat het bepalen van de grenzen niet alleen over aan de onderbuik van de ervaren journalisten, maar stel ze in goed overleg met de makers en de beroepsgroep op. Zorg dat er voldoende expertise over nieuwe ontwikkelingen in de AI is. Versterk de controle op beeldmanipulatie omdat het vaststellen van beeldmanipulatie of beeldgeneratie met het blote oog niet meer volstaat. Vergroot de kennis van digitale beelddetectietools en stel procedures op om te controleren of aangeleverde beelden authentiek zijn. Investeer als redactie in de opleiding van medewerkers om zowel technische kennis van AI-beeldproductie en -analyse te versterken, als bewustzijn van de ethische valkuilen te vergroten. Experimenteer met de nieuwe mogelijkheden, maar stel ook duidelijke grenzen op basis van de richtlijnen. Documenteer en evalueer die experimenten om de journalistieke geloofwaardigheid te waarborgen. Bied transparantie op maat aan het publiek (bijvoorbeeld watermerken, metadatamarkeringen), zodat het weet wat de aard van het journalistieke beeld is en blijft vertrouwen op de journalistieke integriteit van de redactie. Zoek balans in de detaillering van de transparantielabels of indicaties van AI-inzet. Betrek het publiek in de discussie over AI-gebruik: geef informatie, bied een podium voor discussie en laat de ombudspersonen of hoofdredactie het beleid toelichten. Wees kritisch op door AI gereproduceerde stereotyperingen en vooroordelen. Bij AI-beeldgeneratie is een precieze en gedetailleerde prompt essentieel om te voorkomen dat AI stereotypen reproduceert.

Voor een verantwoorde en ethische omgang met AI in journalistieke beeldproductie zijn duidelijke richtlijnen belangrijk. Deze richtlijnen bevatten de volgende elementen:

- Journalisten behouden altijd de eindverantwoordelijkheid en zijn transparant zijn over AI-gebruik, bijvoorbeeld via labels, watermerken of metadata;
- AI-gegenereerde beelden wekken geen misleidende indruk of versterken geen schadelijke stereotypen, en worden vermeden bij feitelijke nieuwsverslaggeving;
- Auteursrechten worden gerespecteerd door alleen AI-tools te gebruiken die op legaal verkregen data zijn getraind.

Ten slotte

AI biedt zowel kansen als risico's voor journalistieke beeldproductie. Het vergt nieuwe vaardigheden, transparantie en ethische reflectie van beeldprofessionals. Het succes van AI in de journalistiek zal ook afhangen van duidelijke richtlijnen, technologische controlemechanismen en een open dialoog met het publiek. Dit onderzoek zet een stap om deze uitdagingen in kaart te brengen en oplossingsrichtingen aan te reiken.

Samenvatting

Samenvatting van de samenvatting

Hoe kan de journalistiek op een verantwoorde manier omgaan met de toenemende mogelijkheden en risico's van beeldbewerking en -productie met AI?

Mogelijkheden

Nieuwe manieren van beeldproductie en tijdbesparing bij onder meer:

- Infographics
- Visualisaties van niet-fotografeerbare zaken
- Satirische beelden
- Illustraties
- Illustratieve foto's

Risico's

- **Reproductie van stereotypen** in gegenereerde afbeeldingen
- **Kwaliteitsdaling:** vervanging van genuanceerd mensgemaakt beeld door AI-eenheidsworst.
- **Schending van beeldrechten** door 'gestolen' trainingsdata.
- **Verlies van de bewijswaarde** van foto's
- **Baanverlies** bij (foto)journalisten

Daarnaast heeft AI-beeldgeneratie een klimaatimpact, en gebeurt de training van de meest bekende AI-modellen door 'data labellers' die werken in slechte arbeidsomstandigheden.

De geïnterviewde experts pleiten naast AI-beeldrichtlijnen op redactieniveau voor **sectorbrede regulering** door beroepsverenigingen, vakbonden en raden voor journalistiek. Beeldmakers, ombudspersonen en redacteuren benadrukken tegelijk het belang van de **journalistieke vrijheid**.

Vijf basisprincipes voor verantwoord journalistiek AI-gebruik*

1. Journalisten **houden het heft in handen** en hebben eindverantwoordelijkheid over beeldproductie (mens-machine-mensprincipe)
2. Journalisten **zijn transparant over AI-gebruik**, maar waken voor een 'informatie-overload'
Hoe uitgebreid en nadrukkelijk die transparantie moet worden doorgevoerd, hangt af van het journalistieke genre en de specifieke beeldtoepassing (zie #4)
3. Journalisten **behoeden zich voor de reproductie van stereotypen** in AI-beeld
Hiervoor is het belangrijk dat redacties investeren in de opleiding van journalisten, zowel op technisch gebied (hoe laat je AI doen wat je wilt) als over beeldvorming en stereotypen.
4. Journalisten maken een duidelijk **onderscheid tussen illustratief beeld en 'feitelijk beeld'**
Om het vertrouwen in de journalistieke foto als 'bewijs' voor echte situaties te beschermen, is het volgens beeldredacteuren belangrijk dat er genrespecifieke normen worden gehanteerd:

Journalistiek genre	Beeldgeneratie	AI-fotobewerking
Satirische beelden	Geen extra beperkingen	Geen extra beperkingen
Niet-fotorealistische illustraties	Geen extra beperkingen	Niet van toepassing
Fotorealistische illustraties	Grote vrijheid, mits extra transparantiemaatregelen om misleiding te voorkomen	Niet van toepassing
Illustratieve foto's	Nooit toegestaan. AI-beelden mogen niet als 'foto's' gepresenteerd worden	Naast belichting en kleurcorrectie mogen in sommige gevallen ook elementen worden weggehaald
Nieuwsfoto's	Nooit toegestaan.	Alleen voor belichting, kleurcorrectie en het weghalen van lensstofjes

5. Journalisten **respecteren de rechten van beeldmakers**
Om verantwoord gebruik te kunnen maken van AI-beeldgeneratietools is het cruciaal dat journalisten zich ervan vergewissen dat het trainingsmateriaal rechtmatig verkregen is.

*op basis van de interviews en het literatuuronderzoek

Hoofdstuk 1

Inleiding

Inleiding

1.1 Aanleiding

Zomer 2023 moest NOS door het stof omdat de nieuwsorganisatie per abuis een nep-foto met talrijke pinguïns had laten zien in het NOS Journaal (NOS, 25 augustus 2023). De redactie had niet in de gaten gehad dat de foto gegenereerd was met AI (artificiële intelligentie of kunstmatige intelligentie). Op 22 mei 2024 ging een foto van een explosie bij het Pentagon viral op sociale media, met als gevolg een daling van de Amerikaanse beurskoersen. Dat was niet nodig geweest, want de foto was nep, er was geen explosie bij het Pentagon (The Guardian). Ook deze foto was gemaakt met AI. Regionale omroepen krijgen veel foto's en filmpjes van burgers toegestuurd ter publicatie. Het is voor de omroepredacteurs in toenemende mate problematisch om – zeker onder tijdsdruk – vast te stellen wat het echtheidsgehalte van dit beeldmateriaal is.

Bovengenoemde voorbeelden maken duidelijk dat de uitdaging voor journalistieke media evident is: hoe dienen ze om te gaan met de toegenomen risico's van met AI gegenereerd beeldmateriaal?

Aan de positieve kant belooft AI ook veel kansen om fraai beeld te produceren in korte tijd, wat bij journalistieke producties een meerwaarde kan hebben. Hoe kunnen journalistieke media op een verantwoorde manier gebruikmaken van al deze nieuwe mogelijkheden?

Recente AI-ontwikkelingen hebben ervoor gezorgd dat iedereen met een computer of smartphone en een internetverbinding fotorealistische beelden kan fabriceren. Mediaorganisaties moeten dus nieuwe omgangsvormen ontwikkelen met beeld, zowel de aan hen aangeboden foto's en video's als zelfgemaakte beelden. De urgentie van nieuwe beeldrichtlijnen lijkt groot. De ontwikkelingen gaan zo snel dat een internationale groep betrokken ontwikkelaars, ondernemers en techpioniers zelfs gepleit heeft voor een moratorium (onderzoekspauze) in het verder ontwikkelen van deze technologie (The New York Times, 23 maart 2023).

AI biedt veel nieuwe mogelijkheden, zoals tijds- en kostenbesparingen, maar brengt ook nieuwe ethische vraagstukken aan het licht. In tegenstelling tot bij de meer 'traditionele' manieren van beeldproductie, is er op het gebied van AI-technologie nog nauwelijks consensus bij makers en publiek over de mogelijkheden en kansen enerzijds en ethische grenzen in het gebruik van deze nieuwe technologie anderzijds. Meer expertise is gewenst bij fotojournalisten en grafisch vormgevers om goed om te gaan met de nieuwe AI-mogelijkheden en bij beeldredacteurs om gemanipuleerd beeld te beoordelen.

Problematisch is daarnaast dat het publieke vertrouwen in de journalistiek mogelijk onder druk komt te staan door de toegenomen mogelijkheden van beeldmanipulatie. Bovendien worden de huidige terughoudendheid en het ongemak van professionele mediaorganisaties om hun beeldbeleid transparant te maken, door deze disruptieve AI-ontwikkelingen versterkt. Zo is er onder meer de clash tussen de objectiviteitsclaims die media in hun beeldrichtlijnen maken versus de praktijk van de fotojournalistiek, waarin bewerkingen en aanpassingen van gemaakt beeld wel degelijk een rol kunnen en soms moeten spelen (Arends & De Jong, 2024). Daarnaast lijkt transparantie over de totstandkoming van beeld niet altijd tot meer vertrouwen in de journalistiek te leiden (Dierickx et al., 2024).

1.2 Onderzoeksvragen

De hoofdvraag van dit onderzoek luidt:

Hoe kan de journalistiek op een verantwoorde manier omgaan met de toenemende mogelijkheden en risico's van beeldbewerking en -productie met AI?

Deelvragen:

1. Welke mogelijkheden biedt AI aan beeldproductie voor journalistieke media?
2. Welke risico's zijn er in het gebruik van AI bij beeldproductie voor journalistieke media?
3. Welke richtlijnen hebben redacties momenteel voor het omgaan met AI-beelden?
4. Welke expertise hebben nieuwsredacties nodig bij het selecteren en produceren van AI-beeld?
5. Welke vormen van transparantie zijn gewenst over de totstandkoming van het journalistieke AI-beeld?

1.3 Methode

Om op bovenstaande onderzoeksvragen een antwoord te geven werd in deskresearch een literatuuronderzoek verricht naar (beeld)richtlijnen in de journalistiek en vertrouwen in de media. Ook is er in die fase een inventarisatie opgesteld van veelgebruikte visuele AI-applicaties.

In dit kwalitatieve onderzoek zijn er twee opeenvolgende series semigestructureerd diepte-interviews afgenomen. In de eerste fase is een diverse groep experts ondervraagd, onder wie gespecialiseerde academici, juristen, belangenvertegenwoordigers, beeldanalisten. Op basis van de resultaten van die eerste interviews is een vragenlijst opgesteld gericht op beeldmakers in de journalistieke praktijk, onder wie fotojournalisten, beeldredacteuren, grafisch vormgevers en ombudspersonen (die laatste groep bestaat niet uit makers, maar mensen die in de praktijk o.a. klachten over journalistieke beelden behandelen).

In totaal zijn er 19 experts en 41 praktijkmensen geïnterviewd. De keuze voor semi-gestructureerd diepte-interviews is ingegeven door de voordelen van deze interactieve benadering, die tot diepgaande inzichten, ontdekking van nuances en groter contextueel begrip leidt. In tegenstelling tot gestandaardiseerde vragenlijsten kunnen onderzoekers zich aanpassen aan de respondenten, waardoor er ruimte is voor onverwachte ontdekkingen en verhelderende details.

Na afloop zijn de interviews getranscribeerd met behulp van GoodTape en AmberScript. Vervolgens codeerden de onderzoekers de transcripten met Atlas.TI, waarbij ideeën zijn geïdentificeerd en gelabeld op basis van specifieke categorieën, thema's en concepten.

1.4 Uitvoering

Het initiatief voor dit onderzoek is genomen door de leerstoelgroep Journalistiek en Nieuwe Media van Universiteit Leiden, onder leiding van prof.dr. Jaap de Jong en dr. Astrid Vandendaele. Samen met onderzoeker Stef Arends, junior onderzoekers Maartje van der Woude en Tijmen Koppelaar, en mediapartners AD (Rennie Rijkma en Derk Westra), KRO-NCRV's Pointer (Jerry Vermanen) en Stichting Regionale Publieke Omroep (Sharla Hagenaars) is dit onderzoek uitgevoerd door de leerstoelgroep Journalistiek en Nieuwe Media, met medewerking van Marc Prüst (Marc Prüst Consultancy).

Het onderzoek is bekostigd door het Stimuleringsfonds voor de Journalistiek (Regeling Onderzoek 2023-2024).

1.5 Opbouw van dit onderzoeksrapport

De opbouw van dit onderzoeksrapport is als volgt: hoofdstuk twee bevat een literatuurstudie over beeldrichtlijnen in de journalistiek, ontwikkelingen op het gebied van AI en beeld, de rol van transparantie en vertrouwen in de media. Hoofdstuk drie behelst een nadere omschrijving van de gebruikte onderzoeksmethoden. Hoofdstuk vier biedt een overzicht van de resultaten uit de diepte-interviews. Hoofdstuk vijf en zes bieden achtereenvolgens conclusies en aanbevelingen gericht op de beroepsgroepen van beeldmakers en beeldbeoordelaars. De bijlage betreft een overzicht van alle geïnterviewden en veelgebruikte visuele AI-applicaties te vinden.

Hoofdstuk 2

Literatuurstudie

Literatuurstudie

De literatuurstudie die volgt beschrijft de rol van beeld in de journalistiek, stelt het begrip publieksperceptie van camera-objectiviteit en realisme ter discussie, beschrijft de relatie tussen vertrouwen in de journalistiek en transparantie in de media en inventariseert mogelijkheden en risico's van beeld AI in de journalistiek, vat de bestaande activiteit samen om AI-beeldrichtlijnen op te stellen en in te voeren en schetst ten slotte ethische dilemma's rond beeld-AI.

2.1 Beeld in journalistiek

Beeld is een essentieel onderdeel in de journalistiek. Hoewel het nieuws lang is gedomineerd door woorden en tekst (Thomson 2019), is het belang van fotografie en video in de journalistiek onomstreden. Visuele elementen zijn cruciale componenten in het huidige medialandschap en de verslaggeving (Grabe & Bucy, 2009; Gritmann 2019; Koltermann 2019; Thomson 2022).

Door de recente technische ontwikkelingen op het gebied van beeld-AI zijn er discussies ontstaan tussen aan de ene kant de vermeende journalistieke objectiviteit en waarheidsgehalte van fotorealisme en aan de andere kant het subjectieve gebruik van journalistieke beelden (Thomson et al. 2024; Arends & De Jong 2024).

2.2 Objectiviteitsclaim nieuwsfoto's

Al voor de recente technologische ontwikkelingen op het gebied van beeld-AI was er een debat over de grote subjectiviteit en manipulatiemogelijkheden in fotografisch beeld. Foto's bieden een realiteitsuggestie die maakt dat het publiek ze vaak beschouwt als bewijs voor echtheid, ondanks de geconstrueerdheid van het beeld (Newton 1998; Messaris & Abrahams 2001, 217). Zelizer (2010) beschrijft hoe fotorealisme ertoe leidt dat het publiek haar ongeloof opschort, en Gynnild (2019) stelt dat live omroeptechnologie de indruk kan wekken dat visuele verslaggeving een ruwe en ongeredigeerde representatie van gebeurtenissen is (Thomson et al. 2024).

In de fotojournalistiek zijn richtlijnen opgesteld voor acceptabele en onacceptabele manipulaties van het beeld. Die richtlijnen zijn vaak eerder opgesteld om voor het publiek de grenzen van serieuze fotojournalistiek af te bakenen ten opzichte van minder serieuze journalistiek (ook wel 'boundary rhetoric' genoemd) dan om reëel inzicht te geven in de subjectiviteit van fotografisch journalistiek beeld (Arends & De Jong 2024).

De nieuwe ontwikkelingen in AI maken de productie van fotorealistische beelden steeds eenvoudiger. Welke mogelijkheden van beeld-AI en kansen in de journalistiek treffen we in de literatuur aan en wat zijn de beschreven risico's?

2.3 Mogelijkheden van beeld-AI in de journalistiek

In de visuele journalistiek biedt kunstmatige intelligentie kansen voor tijd- en kostenbesparing, creatieve brainstorming, het genereren van illustraties, de productie van infographics, datavisualisaties en alternatieve stockbeelden en het visualiseren van toekomstbeelden en zaken die moeilijk of niet te fotograferen zijn (Thomson, Thomas & Matich, 2024, op basis van twintig interviews met beeldredacteuren uit zeven westerse landen, maar niet uit Nederland).

Thomson et al. (2025) maken daarom een onderscheid tussen onzichtbare AI-toepassingen (zoals beeldbewerking, upscaling en automatische tagging) en zichtbare AI-content (volledig gegenereerde afbeeldingen). AI-gedreven technologie stelt journalisten in staat om traditionele beperkingen te overstijgen door ze te faciliteren in een efficiënte en accurate verzameling, analyse en verspreiding van informatie (Dierickx et al. 2024). Deze voordelen kunnen ook betrekking hebben op beeldmateriaal.

2.3.1 Tijd- en kostenbesparing

Enkele jaren geleden was het nog tijdrovend om overtuigende visuals te maken door middel van AI, maar de apps en programma's worden steeds beter en sneller (Hausken, 2024). En afgebeelde personen met zes vingers worden steeds schaarser.

Managers in de media hebben vaak een scherp oog voor de mogelijkheden van kostenbesparingen. Een beeld dat geproduceerd is met een rake prompt kan sneller en goedkoper tot stand komen dan een foto die een professionele fotograaf soms na een tijdrovende reis en onder moeilijke omstandigheden en na diverse fotobewerkingen heeft kunnen maken. Het lijkt kosteneffectiever om een beeldredactie te trainen in de nieuwste beeld-AI-apps dan om kostbare fotografen in te schakelen (Noain-Sánchez, 2021).

2.3.2 Creativiteit in brainstormfase

Generatieve kunstmatige intelligentie kan journalisten bijstaan bij het brainstormen van ideeën (Wu, 2024). Een gegenereerde afbeelding hoeft niet als eindproduct te worden gebruikt om de journalistiek van dienst te kunnen zijn, maar kan dienen als inspiratiemiddel.

2.3.3 Productie van alternatieve stockfoto's, illustraties, infographics en datavisualisaties

Generatieve AI heeft de potentie om de productie van stockfoto's, datavisualisaties en illustraties te ondersteunen, vergemakkelijken of vervangen.

Stockfoto's zijn kant-en-klare beelden die dienen als visueel materiaal voor onder meer journalistieke media (Mortensen et al., 2023). Via stockbureaus kunnen klanten direct foto's kopen en gebruiken. In de journalistiek worden stockfoto's steeds vaker ingezet, maar die sluiten niet aan bij een specifiek nieuwsverhaal. Journalisten hechten juist veel waarde aan die accuratesse en context van een passende afbeelding bij een specifiek nieuwsverhaal (Mortensen et al., 2023). Volgens Thomson, Thomas & Matich (2024) ligt een mogelijke oplossing in gegenereerde afbeeldingen: deze kunnen snel en goedkoop op maat worden gemaakt voor nieuwsverhalen. Voor de vervanging van stockbeeld zou AI een oplossing kunnen bieden.

Datavisualisatie is niet weg te denken uit de hedendaagse journalistiek, omdat het een daadkrachtig communicatiemiddel is om het publiek te informeren (Fu & Stasko, 2024). Een vorm van datavisualisatie is een infographic. Een infographic bestaat uit informatie, kennis, teksten en afbeeldingen die verschijnen in de vorm van een eenvoudig en beknopt visueel ontwerp dat de lezer helpt om ideeën op te nemen, en die doorgaans een grote hoeveelheid informatie in weinig ruimte laten zien (Hamza, 2023; Weber, 2013). Grafisch vormgevers kunnen speciale AI-applicaties gebruiken om snel en daadkrachtig de meest complexe en uiteenlopende infographics te produceren (Taha et al., 2024).

Voor de journalistiek is het gebruik van gegenereerde afbeeldingen om artikelen te illustreren potentieel aantrekkelijk, omdat het informatieve doeleinden kan dienen, geschreven verhalen kan ondersteunen en kan bijdragen aan het visualiseren van gebeurtenissen in het verleden of de toekomst (Schröter, 2023).

2.3.4 Visualisatie van niet-fotografeerbare zaken

Kunstmatige intelligentie leent zich uitstekend voor het genereren van zaken die niet fotografeerbaar zijn (Li et al., 2024). Denk aan verbeeldingen van een hypothetische toekomst, een ver verleden of een specifiek deel van het heelal. Voorheen waren journalistieke media aangewezen op minder adequate stockbeelden of illustraties voor nieuwsverhalen over deze onderwerpen. Nu kunnen zij zelf niet-fotografeerbare zaken genereren, waardoor elk nieuwsverhaal voorzien kan worden van een passend en uniek beeld.

2.4 Risico's van beeld-AI in de journalistiek

De meeste aandacht in de literatuur gaat uit naar de risico's die de nieuwe tools met zich meebrengen voor de journalistiek. Ze hebben betrekking op verliezen van vertrouwen in de journalistiek, op het schenden van beeldrechten van fotografen en media, op baanverlies voor fotojournalisten, en op het ontstaan van kennisilo's op het gebied van deze innovaties in de journalistieke praktijk.

2.4.1 Dalend vertrouwen in de geloofwaardigheid van nieuws

Schandalen over de digitale manipulatie van foto's door fotografen droegen bij aan een toenemend kritische houding van het publiek tegenover fotografie in de journalistiek (Arends & De Jong, 2024), al voordat kunstmatige intelligentie haar opmars maakte. Inmiddels gaan verhitte discussies over de verantwoorde inzet van kunstmatige intelligentie in de nieuwsfotografie hand in hand met debatten over het gevaar van deepfakes, gemanipuleerde foto's en apocalyptische waarschuwingen over een post-truth samenleving (Hausken, 2024). De negatieve invloed van AI-afbeeldingen op vertrouwen in nieuwsorganisaties wordt door journalistieke beeldmakers genoemd als een bedreiging voor de geloofwaardigheid van nieuws en daarmee de journalistieke functie van 'leverancier van de waarheid' (idem).

Bovendien zijn computationele technieken voor het genereren van fotorealistische afbeeldingen inmiddels zo geavanceerd dat fotografen en beeldredacteurs niet in staat blijken om foto's te onderscheiden van computer-gegenereerde afbeeldingen (Lehmuskallio et al. 2019). Ook blijkt zelfs dat mensen synthetische gezichten als geloofwaardiger kunnen beoordelen dan echte (gefotografeerde) gezichten (Nightingale & Farid 2022). Het onderzoek van Thomson et al (2025) wijst uit dat ook veel nieuwsconsumenten niet doorhebben wanneer AI wordt gebruikt in journalistieke beelden en dat AI-gebruik zo ongemerkt wordt geaccepteerd.

Tegelijkertijd zorgt de huidige digitale samenleving ervoor dat iedereen via sociale media visuele boodschappen kan publiceren, waardoor het beoordelen van (visuele) user-generated content tot de dagelijkse routines van een (beeld)redactie behoort (Khan et al. 2022). De beeldredacteur heeft dus als poortwachter een steeds moeilijker taak om echte foto's van AI-gegenereerde beelden te onderscheiden. (Thomson et al. 2024)

2.4.2 Copyrightproblemen

Net als vele anderen concluderen Jones et al. (2023) dat generatieve AI inbreuk kan maken op intellectuele eigendomsrechten en auteursrechtelijk beschermd materiaal verspreiden zonder toestemming, waardoor journalisten en nieuwsorganisaties het risico lopen op juridische claims als ze deze systemen en hun output gebruiken. Het is ook onduidelijk of output die met het gebruik van deze tools wordt gemaakt, auteursrechtelijk beschermd kan zijn. De reden hiervoor is dat modellen getraind worden op massa's gegevens die van het web zijn geschraapt zonder toestemming van de makers van die gegevens. In hun rapport bespreken Thomson, Thomas, Riedlinger en Matich (2025) de auteursrechtelijke implicaties van AI-gegenereerde beelden, met zorgen over de niet-transparante training van AI-modellen, het ontbreken van bronvermelding en mogelijke auteursrechtsschendingen. AI-tools kunnen intellectueel eigendom van kunstenaars en fotografen gebruiken zonder compensatie, terwijl fact-checking en creditering bemoeilijkt worden. In de journalistiek, aldus Thomson et al (2025) leidt AI-gebruik tot risico's op plagiaat en verlies van inkomsten voor originele bronnen. Juridische gevolgen hiervoor blijven voorlopig nog onzeker, en regelgeving evolueert nog.

2.4.3 AI kan vooroordelen versterken

Uit onderzoek blijkt dat met AI gegenereerd beeld vaak een bias bevat en vooroordelen bevestigt. Een van de risico's van het gebruik van zulke beelden is de verspreiding van schadelijke vooroordelen en stereotypen. Een van de meest gebruikte generatoren, Stable Diffusion, reproduceert aantoonbare raciale vooroordelen en genderstereotypen (Cauhan et al. 2024). In het algemeen vonden Cauhan et al. dat afbeeldingen van banen in de horeca of van laagbetaalde banen werden gedomineerd door vrouwen, waardoor het stereotype dat vrouwen meer geschikt zijn voor huishoudelijk werk werd bevestigd. Ook Thomson et al. (2025) tonen aan dat AI-systemen inherente biases kunnen bevatten. Hoewel AI bekende biases vertoont tegen bijvoorbeeld vrouwen en mensen met een migratieachtergrond, heeft hun onderzoek ook minder bekende biases geïdentificeerd.

Dit omvat omgevingsbias (waarbij AI de voorkeur geeft aan stedelijke omgevingen boven niet-stedelijke wanneer geen specifieke context is opgegeven in een prompt), rol- en vaardigheidsbias (waarbij vrouwen minder vaak worden afgebeeld in gespecialiseerde functies en mensen met een beperking grotendeels worden genegeerd) en klassebias (waarbij personen uit de ogenschijnlijk middenklasse, oververtegenwoordigd zijn).

De risico's van stereotyperende beelden stammen natuurlijk al van voor het AI-tijdperk. Niet alleen beeld-AI kan bias bevatten, maar uit zowel Nederlands als internationaal onderzoek blijkt dat menselijke mediaberichtgeving ook gender- en racistische stereotypes produceert. Die benadeelt onder meer moslims, mensen op de vlucht en mensen met een migratieachtergrond in het algemeen (Ashmelash & Remmers, 2014; Balçık, 2019; Petca et al., 2012; Georgiou & Zaborowski, 2017).

2.4.4 Vrees voor baanverlies

De fotojournalistiek is al lang onderhevig aan verandering door opkomende technologie. Tussen 1999 en 2015 hebben talloze mediaorganisaties zich geheel ontdaan van hun fotopersoneel, of deze groep drastisch laten inkrimpen (Thomson, 2018). Mogelijk baanverlies van journalisten blijft een heet hangijzer in gesprekken over AI-integratie in journalistieke organisaties (Gutiérrez-Caneda et al, 2024). Zorgen bestaan niet alleen uit complete vervanging van journalisten door AI, maar ook uit gedeeltelijk baanverlies door de overname van routineuze taken door kunstmatige intelligentie (idem). Ook Thomson et al. (2025) bespreken de impact van AI op werkgelegenheid in journalistiek, waarbij vooral fotojournalisten en beeldredacteuren kwetsbaar zijn. De perceptie van AI als een 'goedkope vervanger' van mensen wordt volgens de onderzoekers vooral gevoed door managers, terwijl journalisten zelf AI vaak als een ondersteunend hulpmiddel zien. De opmars van generatieve, visuele kunstmatige intelligentie voedt onder de beroepsgroep de angst dat managers en bedrijfseigenaren de technologie gebruiken om productiviteit en winstmarges te optimaliseren, waardoor creatieve menselijke arbeid overbodig raakt (Boyles en Meisinger, 2020, Bender, 2024; Matich et al., 2024).

2.4.5 Beperkte AI-expertise en kennisilo's

Dodds, Vandendaele et al. (2024) onderzochten hoe kennisilo's – interne barrières voor informatie-uitwisseling binnen organisaties – de adoptie van verantwoorde AI-praktijken belemmeren. Uit 14 semi-gestructureerde interviews met redacteuren, journalisten en managers van de Telegraaf, de Volkskrant, NOS en RTL Nederland blijkt dat gebrekkige kennisdeling en afhankelijkheid van externe AI-leveranciers de implementatie van AI in de journalistiek bemoeilijken.

De studie identificeert vier soorten kennisilo's: verticale (tussen hiërarchische lagen), horizontale (tussen afdelingen), interne (binnen organisaties) en externe (tussen redacties en moederbedrijven).

Het beleidsadvies van onderzoekers op de AI Policy Research Summit (2024) benadrukt eveneens het belang van multidisciplinaire samenwerking en capaciteitsopbouw door AI-training binnen redacties. Journalistieke redacties missen immers vaak de technische expertise om AI-systemen kritisch te beoordelen, en zonder structurele maatregelen riskeren nieuwsorganisaties het verlies van controle over AI-gebruik en journalistieke normen.

2.4.6 Klimaatimpact

Het rapport voortgekomen uit de AI Policy Research Summit in Stockholm (2024) benoemt het belang van onderzoek naar verantwoord AI-beleid en -beheer, met de nadruk op ethische overwegingen, duurzaamheid en inclusiviteit. De onderzoekers waarschuwen voor de ecologische impact van AI, zoals hoge energieconsumptie en e-waste en benadrukt de verantwoordelijkheid van mediabedrijven om duurzamer met AI om te gaan, bijvoorbeeld door energie-efficiënte modellen te gebruiken zoals (de Chinese ontwikkelaars van DeepSeek claimen) en datacentra met groene stroom te selecteren.

Naar de exacte klimaatimpact van de verschillende AI-beeldgeneratieprogramma's is nog weinig onderzoek gedaan, maar uit een recente studie van het Allen Institute for AI (Lucioni et al., 2024) bleek wel dat het genereren van nieuwe afbeeldingen met behulp van AI veel vervuilerend is dan andere toepassingen van kunstmatige intelligentie, zoals het schrijven en samenvatten van teksten of het detecteren van objecten in een afbeelding.

2.5 Journalistieke normen

De journalistiek is geen beschermd beroep, maar kent wel normen en waarden, vastgelegd in redactiestatuten en stijlboeken, die het dagelijkse werk van journalisten sturen. Onderzoekers identificeerden verschillende kernwaarden die essentieel zijn in het journalistieke werkproces waaronder waarheid en geloofwaardigheid, onpartijdigheid en democratische verantwoordelijkheid, originaliteit en creativiteit, integriteit en eerlijkheid, transparantie, onafhankelijkheid, en diversiteit en inclusie (Komatsu et al., 2020; Helberger et al, 2022; Oh en Jung, 2024). Komatsu en collega's (2020) stellen dat AI-systemen niet alleen journalistieke waarden moeten ondersteunen, maar deze ook moeten belichamen. Dit betekent dat ontwerpers rekening moeten houden met hoe AI transparantie, verantwoordelijkheid en neutraliteit kan waarborgen in journalistiek werk. Diakopoulos & Cools (2023) wijzen daarnaast op het belang van aansprakelijkheid en privacybescherming.

2.6 Regelgeving

Binnen de EU is in 2023 de eerste AI-wet ter wereld aangenomen (O'Carroll 2023). Die wet reguleert de toepassing van AI in verschillende industrieën, maar heeft geen betrekking op de journalistiek. Het journalistieke domein wordt geacht om door middel van zelfregulatie met AI om te gaan.

Journalisten en redacteurs baseren werkgerelateerde beslissingen in de praktijk vaak op hun ‘onderbuikgevoel’ en laten hun nieuwskeuzes daarmee vanzelfsprekend en zelfverklarend lijken (Schultz, 2007). Het ‘onderbuikgevoel’ wordt ook geraadpleegd bij keuzes voor het wel of niet gebruiken van generatieve kunstmatige intelligentie in de journalistiek (Diakopoulos & Cools, 2024).

Verlies van publieksvertrouwen en risico's ten gevolge van copyrightscheidingen zijn echter te ernstig om AI-beeldgebruik aan de onderbuik over te laten. Maar hoe dan te opereren.

In het rapport gepubliceerd na de AI Policy Research Summit in Stockholm (2024) stelt men dat AI-governance niet alleen door overheden, maar ook door zelfregulerende initiatieven binnen sectoren moet worden gedragen. Dit is een relevant argument in de journalistieke context, waar zelfregulering en ethische richtlijnen vaak effectiever zijn dan wettelijke beperkingen die persvrijheid kunnen inperken.

Harde, specifieke regulering kan journalistiek werk dus belemmeren en te losse regulering dient geen nut (Gutiérrez-Caneda et al. 2024). Onderzoek toont aan dat zelfregulering en de uitbreiding van ethische principes en richtlijnen binnen redacties werkbaar opties zijn, onder meer omdat regulering van buitenaf botst met kenmerken die inherent zijn aan de journalistiek. (Porlezza, 2023). Het ontwikkelen van richtlijnen die specifieke betrekking hebben op de omgang met kunstmatige intelligentie kan positief zijn omdat het nieuwsorganisaties in staat stelt om te reflecteren op ethische codes en de plek van kunstmatige intelligentie in die codes. Daarnaast zijn AI-richtlijnen nuttig voor het informeren van het publiek door transparantie over keuzes en gewoonten van redacties te bieden (Gutiérrez-Caneda et al. 2024).

Veel mediabedrijven hanteren richtlijnen, maar van standaardisatie is geen sprake (Gutiérrez-Caneda, Lindén & Vázquez-Herrero, 2024). Huidige journalistieke AI-richtlijnen gaan onder meer over transparantie, verantwoordelijkheid, menselijk toezicht en ethische standaarden (Porlezza et al. 2024; Becker et al. 2023; de-Lima-Santos et al. 2024).

De Amerikaanse News Media Alliance en de Deutscher Journalisten-Verband waren in april 2023 de eerste belangengroepen die AI-richtlijnen hebben opgesteld. Tot de vroegste AI-beeldrichtlijnen horen die van tijdschrift WIRED in juni 2023, waarin het volgende staat: ‘personeel mag AI-instrumenten gebruiken om visuele ideeën op te wekken’ en ook “personeel mag AI-gegenereerde afbeeldingen en video's publiceren, weliswaar met bekendmaking van de makers die generatieve AI als instrument gebruiken.”

WIRED is ook duidelijk over wat niet mag: “WIRED zal geen AI-beelden in plaats van stockfotografie gebruiken.” Deutsche Welle publiceerde richtlijnen in september 2023, die medewerkers instrueren om AI niet te gebruiken voor de productie van fotorealistische beelden, maar wel om illustraties en datavisualisaties te creëren of verbeteren (Kasper-Claridge 2023, Thomson et al. 2024).

Er is dus ruimte voor behoedzaam experimenteren, maar de inzet van beeld-AI bij fotografie in de nieuwsjournalistiek is zeer problematisch (Cools en Diakopoulos 2023). Inzet bij ondersteunende illustraties zijn wel mogelijk, maar wel geldt dan het zogenoemde Human in the loop-principle. ANP formuleert het haar AI-leidraad (2023) als volgt:

“We houden in onze productieketen vast aan de nu al geldende lijn mens> machine>mens waarbij het bedenken en beslissen bij de mens begint en eindigt, mogelijk in de tussenliggende productie geassisteerd door de computer.”

Uit recent onderzoek is gebleken dat nieuwsmakers weliswaar experimenteren met AI, maar dat ze zich meestal verzetten tegen fotorealistische tekst-naar-beeld-generatie omdat dit de representatieve integriteit van de journalistiek schaadt (Thomson, Thomas, Matich, 2024).

2.7 Transparantie

AI-gebruik in de journalistiek is vaak onduidelijk voor het publiek (Gutiérrez-Caneda, Lindén & Vázquez-Herrero 2024). De Europese AI Act stelt dat er transparantie vereist is wanneer AI wordt gebruikt voor contentcreatie, zowel beeld als tekst. Nieuwsorganisaties zijn echter vrijgesteld van veel van deze vereisten. Met andere woorden: journalistieke redacties zijn niet verplicht om AI-gebruik expliciet te vermelden, tenzij het om audiovisuele deepfakes gaat of om content met een duidelijk publiek belang. Bovendien worden nieuwsorganisaties niet geclassificeerd als een “hoge-risico” sector, waardoor er minder stringente regelgeving op hen van toepassing is.

Veel studies beschrijven de inzet van AI-beeld in de journalistiek als een bedreiging van het vertrouwen dat het publiek heeft in de journalistiek (bronnen). Daarom blijken veel media beeld-AI in nieuwsdomein te verbieden. Dat verbod wordt gecommuniceerd met het publiek in de stijlboeken en richtlijnen van die media (zie 2.7). Voor media die het beeld-AI-kind niet met het badwater willen weggooien, is er de oplossing van transparantie. Uit het onderzoek van Thomson et al (2025) blijkt dat zowel journalisten als publiek transparantie over AI-gebruik inderdaad belangrijk vinden, maar dat er geen consensus is over hoe dit het best kan worden gecommuniceerd. Vaak wordt geopperd dat de journalistiek transparant moet zijn wanneer, hoe en waarom bepaalde vormen van AI in beeld zijn ingezet. Dat kan met uitgebreide beschrijvingen, zelfs met de gebruikte prompt die tot het geproduceerde beeld heeft geleid. NU.nl meldt het gebruik van een beknoptere vermelding, in de vorm van emblemen:

Uit eerder onderzoek blijkt dat het publiek gemengde gevoelens heeft over AI-gebruik in journalistiek. Sommigen beschouwen AI-gegenereerde nieuwsberichten als objectiever en feitelijker omdat AI geen emotionele of ideologische voorkeuren heeft. Anderen wantrouwen AI-nieuws vanwege de complexiteit van de technologie en de onzichtbaarheid van de onderliggende besluitvormingsprocessen. Dit wordt de “disclosure paradox” genoemd: hoewel het publiek transparantie over AI-gebruik eist, kan expliciete vermelding van AI zelfs leiden tot een afname van vertrouwen in de inhoud (Altay & Gilardi, 2024).

Onderzoek naar de effecten van het labelen van AI-gegenereerd nieuws is nog gaande, maar de eerste resultaten suggereren dat de invloed op publieksvertrouwen negatief is (Toff en Simon, 2023). Zelfs als de nieuwsconsument zo'n label wenst, blijkt dat zij nieuws dat is gelabeld als AI-gegenereerd als minder betrouwbaar acht. Transparantie kan dus niet gezien worden als Haarlemmerwonderolie voor het terugwinnen van vertrouwen.

2.8 AI-tools in de journalistiek

Sinds de meest recente AI-boom in 2022 op gang kwam, experimenteren mediaorganisaties volop met het implementeren van generatieve AI in de journalistiek. AI-tools worden steeds breder ingezet voor zowel de productie van tekst en beeld als voor de analyse en verificatie ervan. Volgens Thomson et al. (2025) hebben veel journalisten dusver geen systematische methodes om AI-gegenereerd beeld te detecteren, waardoor ze kwetsbaar zijn voor misinformatie. De razendsnelle technologische ontwikkelingen zorgen dus voor een reeks aan ethische dilemma's (zie resultaten), maar maken het ook moeilijk voor (beeld)redacteuren en fotografen om up-to-date te blijven met de beschikbare software en te begrijpen wat die precies doet.

Niet alle AI-beeldsoftware werkt op dezelfde manier, en veel van de meest gebruikte programma's zijn verre van transparant over de manier waarop ze beelden genereren, bewerken en analyseren. In dit hoofdstuk geven we een overzicht van de belangrijkste op dit moment (februari 2025) beschikbare programma's voor generatieve beeld-AI, en de publiek beschikbare kennis van hun werking.

2.8.1 Technologie in een stroomversnelling

Hoewel de eerste AI-programma's voor beeldgeneratie al in de jaren '60 werden ontwikkeld, is de functionaliteit en toegankelijkheid ervan lang erg beperkt gebleven. De afgelopen tien jaar hebben verschillende technologische vernieuwingen voor een stroomversnelling gezorgd in de ontwikkeling van generatieve beeld-AI, waaronder deep learning (een manier om computers te trainen met algoritmes die geïnspireerd zijn op de werking van het menselijk brein), generative adversarial networks of GAN's (waarbij een beeldherkenningsprogramma een beeldgeneratieprogramma traint om afbeeldingen te genereren die lijken op echte beelden) en diffusion models (die afbeeldingen stap voor stap afbreken tot een ruisbeeld en op die manier leren hoe ze door dat proces om te draaien uit ruis een nieuw beeld kunnen maken).

Het eerste text-to-image beeldgeneratiemodel – waarbij een prompt van de gebruiker wordt vertaald in een afbeelding – werd in 2015 aan de Universiteit van Toronto ontwikkeld door de toen 19-jarige ontwikkelaar Elman Mansimov. Dat model, AlignDraw, was nog zeer rudimentair en kon slechts afbeeldingen van 32 bij 32 pixels genereren, die verre van fotorealistisch waren. De opkomst van deepfake-video's bracht daar een paar jaar later verandering in. In 2017 creëerden onderzoekers aan de Universiteit van Washington een beeldmodel op basis van speeches van ex-president Barack Obama, dat voor veel mensen niet van echt te onderscheiden was.

Filmregisseur Jordan Peele en mediaplatform BuzzFeed waarschuwden in 2018 voor de mogelijke maatschappelijke consequenties van generatieve beeld-AI-software. Peele gebruikte de inmiddels voor iedereen beschikbare software FakeApp om een deepfake-video van Obama te genereren met de kenmerkende BuzzFeed-titel *You Won't Believe What Obama Says In This Video*.

“This is a dangerous time”, waarschuwde de AI-Obama. “Moving forward, we need to be more vigilant with what we trust from the internet. It's a time we need to rely on trusted news sources. It may sound basic, but how we move forward in the age of information is going to be the difference between whether we survive, or whether we become some fucked-up dystopia.” De video ging de wereld rond maakte het gevaar van AI-beelden voor de publieke informatievoorziening voelbaar.

De website *thispersondoesnotexist.com* schudde de journalistiek een jaar later op met niet van echt te onderscheiden ‘foto’s’ van menselijke gezichten. Verschillende Amerikaanse, Israëlische en Saoedi-Arabisch nieuwswebsites werden misleid en publiceerden ingezonden opiniestukken van niet-bestaande Midden-Oostenexperts, compleet met nep-portretfoto en gelinkte nepprofielen op Twitter en LinkedIn. De verschillende nep-experts pleitten voor strengere sancties tegen Iran, en voor nauwere banden met de Verenigde Arabische Emiraten.

In 2021 publiceerde *The Correspondent* een deepfake-video van Mark Rutte die pleit voor een ambitieus klimaatbeleid. Opvallend is dat deze deepfake, in tegenstelling tot die van Jordan Peele drie jaar eerder, niet bedoeld was om bewustwording te creëren over misleiding door AI. Het doel, zo schreef hoofdredacteur Rob Wijnberg in een journalistieke verantwoording, was ‘om te laten zien hoe klimaatleiderschap eruit kan zien’. Volgens Wijnberg ging het hier niet om nepnieuws, maar juist om het tegenovergestelde:

“In plaats van onwaarheden te presenteren alsof ze waar zijn, in een vorm die pretendeert ‘echt’ te zijn, presenteren wij feiten over klimaatverandering en noodzakelijke maatregelen in een vorm waarvan we expliciet vermelden dat deze nep is.”

The Correspondent oogstte lof voor hun gedurfde werkwijze, maar kreeg ook veel kritiek. Zo was niet iedereen overtuigd dat het mediaplatform voldoende maatregelen had genomen om misleiding te voorkomen. Nep-Mark Rutte geeft op 5:40 van de zeven minuten durende video weliswaar expliciet aan dat de video ‘niet echt’ is, maar critici wijzen erop dat dat gedeelte van de video in de vele clips die op sociale media de ronde doen, meestal weggeknipt is.

De ontwikkeling van fotorealistische text-to-image AI-modellen kwam pas in 2022 op gang met de introductie van Open AI's DALL-E 2, en de lancering van concurrenten Midjourney en Stable Diffusion datzelfde jaar. Sindsdien is het gebruik van generatieve beeld-AI in de journalistiek sterk toegenomen. In juni 2022 publiceerden *The Economist* en *Cosmopolitan* de eerste AI-gegenereerde magazine-covers, gemaakt met respectievelijk de modellen Midjourney en DALL-E 2.

De afgelopen drie jaar is het aantal beschikbare text-to-image beeldgeneratieprogramma's geëxplodeerd. Tegelijk wordt beeld-AI ook op steeds meer andere manieren ingezet door redacties, bijvoorbeeld om AI-beelden te detecteren, om beschrijvingen te maken van foto's en om de oorsprong van foto's te achterhalen. Op dit moment is het zo goed als onmogelijk om een compleet overzicht te geven van de beschikbare tools voor beeld-AI in de journalistiek. Hieronder bespreken we daarom een selectie van enkele generatieve AI-modellen die in de journalistiek het meest gebruikt worden, evenals enkele modellen met unieke functionaliteiten. Een uitgebreider overzicht staat in Bijlage 3.

2.8.2 Beeldgeneratieprogramma's

Midjourney - Midjourney Inc.

Functionaliteit: Genereren van beelden op basis van tekst-prompts.

Trainingsdata: Niet transparant. Midjourney geeft geen informatie over welke trainingsdata het gebruikt. Wel is bekend dat het model onder meer getraind is op basis van de LAION-5B-database, een verzameling van zo'n twee miljard beelden en beschrijvingen die automatisch en zonder toestemming van de makers werden gedownload. Zo kwamen onder meer beelden van kindermisbruik en foto's van genocides zoals de holocaust terecht in de data die gebruikt werd om het model te trainen.

Model: Een eigen model, waarvan de broncode niet publiek is.

Bias: Het is aangetoond dat Midjourney bestaande ongelijkheden in de samenleving reproduceert en versterkt. Midjourney Inc. heeft geen duidelijkheid gegeven of en hoe het de reproductie van ongewenste vooroordelen tegengaat. Ook kan het model makkelijk gebruikt worden voor het maken van racistische, gewelddadige en haatdragende beelden, inclusief beelden van bestaande personen.

Vergoeding: Geen vergoeding voor de makers van trainingsbeelden.

Kosten voor de gebruiker: Gratis trial voor 25 afbeeldingen. Daarna 8 dollar per maand voor bedrijven met een jaaromzet van minder dan een miljoen dollar. Anders 60 dollar per maand.

Stable Diffusion van Stability AI

Functionaliteit: Genereren van beelden op basis van tekst-prompts.

Trainingsdata: Onder andere de LAION-5B-database, een verzameling van zo'n twee miljard beelden en beschrijvingen die automatisch en zonder toestemming van de makers werden gedownload. Zo kwamen onder meer beelden van kindermisbruik en foto's van genocides zoals de holocaust terecht in de data die gebruikt werd om het model te trainen.

Model: Een eigen model, waarvan de broncode publiek is.

Bias: Het is aangetoond dat Stable Diffusion bestaande ongelijkheden in de samenleving reproduceert en versterkt, waaronder een algemene overrepresentatie van witte mannen in de gegenereerde afbeeldingen, seksualisering van vrouwen van kleur en stereotyperende afbeelding van inheemse groepen. Ex-CEO van Stability AI gaf aan dat het de 'verantwoordelijkheid van de gebruiker' is om Stable Diffusion op een ethische, morele en legale manier te gebruiken. Het bedrijf zelf legt geen verplichte contentbeperkingen op.

Vergoeding: Geen vergoeding voor de makers van trainingsbeelden.

Kosten voor de gebruiker: Gratis wanneer men het model downloadt en lokaal draait. Ook mogelijk via een API en in combinatie met een clouddienst als Replicate, voor ±1 tot 10 dollarcent per afbeelding, afhankelijk van het specifieke model.

Shutterstock AI Image Generator van Shutterstock Inc.

Functionaliteit: Genereren van beelden op basis van tekst-prompts.

Trainingsdata: De eigen fotodatabank van Shutterstock.

Model: Niet transparant. Shutterstock geeft enkel aan dat het model werd ontwikkeld 'op basis van de beste technologie van onze partners' en 'in strategische partnerschappen met (±) OpenAI, Meta en LG AI Research', maar geeft geen specifieke informatie over welke modellen er precies gebruikt werden, en hoe die werden getraind.

Bias: Shutterstock CEO Paul Hennessy gaf in een interview aan dat de Shutterstock AI image generator is getraind op basis van 'een beeldbank die de diversiteit weerspiegelt van de wereld waarin we leven'. Er wordt echter geen specifieke informatie gegeven over hoe het bedrijf omgaat met het risico op stereotypering en versterking van vooroordelen in beelden.

Vergoeding: Vergoeding voor de makers van trainingsbeelden. Shutterstock keert de makers een niet-gespecificeerd percentage van de verkoopprijs van de AI-beelden uit via een Contributor Fund.

Kosten voor de gebruiker: €6 per maand, voor 100 beeldgeneraties per maand.

Dall-E 3 van OpenAI

Functionaliteit: Genereren van beelden op basis van tekst-prompts en bewerken van afbeeldingen, geïntegreerd in ChatGPT

Trainingsdata: Niet transparant. OpenAI geeft aan dat Dall-E getraind is op een ‘combinatie van publiek beschikbare bronnen en bronnen waarvan we een licentie hebben’, maar specificeert niet om welke. Het bedrijf heeft een opt-outformulier geïntroduceerd, waarmee fotografen en artiesten kunnen vragen om hun werk niet te gebruiken voor training.

Model: Een eigen, door OpenAI ontwikkeld model.

Bias: Net als de andere AI-beeldmodellen gebaseerd op grote hoeveelheden trainingsdata, reproduceert en versterkt Dall-E bestaande ongelijkheden in de samenleving. OpenAI geeft aan actief bias tegen te gaan in de Dall-E-modellen. Vrouwen werden in de gegenereerde afbeeldingen vaker geseksualiseerd dan mannen. De maatregelen die het bedrijf nam om dit probleem tegen te gaan, zorgden er echter voor dat vervolgens bij een genderneutraal prompt vaker mannen dan vrouwen werden weergegeven.

Vergoeding: Geen vergoeding voor de makers van de trainingsbeelden. Om copyright-claims tegen te gaan van artiesten wier werk gebruikt werd in de trainingsdata, heeft Dall-E 3 een contentbeperking opgelegd aan gebruikers. Prompts waarin gevraagd wordt om een afbeelding te maken in de stijl van een bestaande artiest, worden hierdoor soms geweigerd, maar niet altijd.

Kosten voor de gebruiker: Gratis als onderdeel van ChatGPT, ook te gebruiken via Microsoft Copilot en Bing Image Creator.

Ideogram AI van Ideogram

Functionaliteit: Genereren van beelden op basis van tekst-prompts. Specifiek geschikt voor beelden waar tekst in voorkomt.

Trainingsdata: Niet transparant. Op de website van Ideogram is niet terug te vinden welke data gebruikt werd voor het trainen van het model.

Model: Een zelf ontwikkeld model.

Bias: Het is aannemelijk dat aangezien Ideogram AI getraind is op grote hoeveelheden van internet gedownloade afbeeldingen, ook Ideogram bestaande ongelijkheden in de maatschappij reproduceert en versterkt. CEO Mohammad Norouzi gaf in een interview aan dat Ideogram ‘bepaalde soorten afbeeldingen heeft verwijderd uit de trainingsdata’, maar dat er ‘meer onderzoek nodig is’.

Vergoeding: Geen vergoeding voor de makers van trainingsbeelden.

Kosten voor de gebruiker: Gratis beperkt tot 40 afbeeldingen per week, anders 8 dollar

 Adobe Firefly van **Adobe Inc.**

Functionaliteit: Genereren van beelden op basis van tekst-prompts.

Trainingsdata: Copyright-vrije beelden uit de beeldbanken van Wikimedia en Flickr, en de eigen fotodatabank van Adobe Stock images. Echter, vorig jaar bleek dat in die ‘eigen’ beeldbank ook door andere AI-programma’s gegenereerde beelden staan, waarvan niet duidelijk is op welke data ze getraind zijn. Zo zou ongeveer 5% van de beeldbank van Adobe Stock bestaan uit met Midjourney gemaakte beelden.

Model en trainingsmethode: Niet transparant. Firefly is gebaseerd op een eigen model genaamd Adobe Sensei.

Bias: Net als de andere AI-beeldmodellen gebaseerd op grote hoeveelheden trainingsdata, reproduceert en versterkt Adobe Firefly bestaande ongelijkheden in de samenleving. Adobe benadrukt wel te proberen dit probleem te verminderen, en communiceert op zijn website over de verschillende manieren waarop.

Vergoeding: Geen vergoeding voor de makers van tr Volgens Adobe hebben de makers van de beelden uit de Adobe Stock-beeldbank hun intellectuele eigendomsrecht verkocht aan het bedrijf, en betekent dit dat het die ook mag gebruiken voor de creatie van nieuwe beelden. Veel van hen zijn het daar echter niet mee eens, en vinden dat er misbruik is gemaakt van hun werk.

Kosten voor de gebruiker: €11,17 per maand of inbegrepen in Adobe Creative Cloud, onbeperkt aantal beeldgeneraties per maand.

 Tess

Functionaliteit: Genereren van beelden op basis van tekst-prompts. Naar eigen zeggen ‘de eerste beeldgenerator die artiesten correct betaalt’.

Trainingsdata: Tess werkt nauw samen met specifieke artiesten die een relatief kleine dataset uploaden en gebruikers in staat stellen om zo afbeeldingen te maken in ‘hun stijl’.

Model: De basis van het model maakt gebruik van Stable Diffusion, dat op zijn beurt weer getraind is met de niet-copyrightvrije database LAION-5B.

Bias: Het is aannemelijk dat aangezien het model gebaseerd is op Stable Diffusion, dat getraind is op grote hoeveelheden van internet gedownloade afbeeldingen en aantoonbaar maatschappelijke ongelijkheden reproduceert, ook Tess een probleem heeft met bias in de gegenereerde afbeeldingen.

Vergoeding: Ja, een uniek vergoedingsmodel waarbij kunstenaars zelf de controle over het gebruik van geüploade beelden behouden.

Kosten voor de gebruiker: Vanaf \$20 voor 50 gegenereerde afbeeldingen per maand.

Grok Image Generation van xAI

Functionaliteit: Genereren van beelden op basis van tekst-prompts en op basis van andere afbeeldingen.

Trainingsdata: Niet transparant. xAI zelf geeft aan dat Grok getraind is op basis van ‘het internet’ en door middel van de ‘AI Tutors’ (medewerkers die instaan voor de training van AI) van het bedrijf. Wel is bekend dat moederbedrijf X de posts op het socialemediaplatform heeft gebruikt voor het trainen van Grok.

Model: Tot voor kort maakte Grok gebruik van het Flux-beeldgeneratiemodel van ontwikkelaar Black Forest Labs, maar inmiddels geeft xAI aan overgestapt te zijn naar een zelf ontwikkeld model genaamd Aurora.

Bias: Net als de andere AI-beeldmodellen gebaseerd op grote hoeveelheden trainingsdata, reproduceert en versterkt Grok bestaande ongelijkheden in de samenleving. Het feit dat eigenaar Elon Musk Grok profileert als een ‘niet-woke’ AI-model, impliceert dat xAI niet bewust bezig is om stereotypering en bias tegen te gaan. Het bedrijf benadrukt bovendien dat het zo min mogelijk restricties oplegt aan de beeldgeneratie met Grok.

Vergoeding: Geen vergoeding voor de makers van trainingsbeelden. Dit terwijl het namaken van bestaande kunstwerken, of het creëren van beelden in de stijl van bekende fotografen met de chatbot ongelimiteerd mogelijk is.

Kosten voor de gebruiker: Gratis beschikbaar via socialemediaplatform X.

2.8.3 Beeldverificatieprogramma's

Naast AI-beeldgeneratie zijn er ook steeds meer AI-tools beschikbaar voor beeldverificatie. Enkele van de meestgebruikte zijn Hive Moderation en AiOrNot. Ook de Universiteit Gent ontwikkelde in samenwerking met de Belgische onderzoeksinstituut Imec en journalisten een AI-beelddetectieplatform dat specifiek gericht is op journalisten. De tool, genaamd Com-Press, geeft meer inzicht in hoe de beelden geanalyseerd worden dan commerciële toepassingen als AiOrNot. Daardoor is echter ook meer kennis vereist om de resultaten goed te kunnen interpreteren dan de commerciële toepassingen die een eenduidiger maar niet altijd correcte waarschijnlijkheidsscore geven of een beeld met AI gemaakt is of niet. De onderzoekers van de UGent en Imec gaven aan dat de tool niet geschikt is voor zelfstandig gebruik door journalisten, maar in samenwerking met AI-experts gebruikt moet worden voor een correcte interpretatie. Hiernaast bespreken we drie detectieprogramma's, en de belangrijkste functionaliteiten en verschillen.

AIORNOT - AIORNOT Inc.

Functionaliteit: Gebruiksvriendelijk online platform voor detectie van AI-gegenereerd beeld. Hoewel niet 100% accuraat, bleek uit tests van fotografiewebsite PetaPixel dat de tool relatief goed onderscheid kan maken tussen AI- en mensgemaakte afbeeldingen.

Model: AIORNOT werkt volgens ontwikkelaar AIORNOT Inc. op basis van een model genaamd Optic.

Kosten voor de gebruiker: \$5 per maand voor 100 maandelijkse beeldcontroles.

Com-Press - Universiteit Gent en IMEC

Functionaliteit: Beeldanalyse op basis van verschillende parameters, zoals de ‘compression fingerprint’, het patroon van ruis in de afbeelding en de metadata. Het platform vergt veel specialistische kennis om correct te interpreteren.

Model: Een combinatie van verschillende beschikbare en zelf ontwikkelde detectiemodellen.

Kosten voor de gebruiker: Gratis beschikbaar.

Truepic Vision - Truepic Inc.

Functionaliteit: Een applicatie die fotografen in staat stelt om een bewijs van authenticiteit toe te voegen aan hun foto's. Op basis van metadata zoals datum, tijd, locatie waar de foto genomen is en de pixels in een afbeelding, creëert Truepic een cryptografische ‘vingerafdruk’ van de foto. Dat is een code of ‘hash’ die wordt opgeslagen in een blockchain, en die mensen die de foto later bekijken, in staat stelt om na te gaan of die is nabewerkt.

Model: n.v.t.

Kosten voor de gebruiker: Op aanvraag.

2.8.4 Beeldherkenningsprogramma's

Een andere door journalisten veelgebruikte toepassing van beeld-AI-software zijn beeldherkenningsstools. Die kunnen worden gebruikt om na te gaan waar een foto voor het eerst geüpload werd, wat kan helpen bij het zoeken van de oorsprong ervan. Maar reverse image search kan ook nuttig zijn wanneer niet duidelijk is wat er precies op een foto te zien is, of waar een foto genomen is. Hierna geven we de drie meestgebruikte voorbeelden:

Tineye - Idée, Inc.

Functionaliteit: everse image search. Door een afbeelding te uploaden, kan men de software laten zoeken naar gelijkaardige afbeeldingen, en websites waar die afbeeldingen voorkomen.

Kosten voor de gebruiker: Gratis beschikbaar

Google Lens - Alphabet, Inc.

Functionaliteit: Reverse image search, maar met de optie om een gedeelte van de afbeelding te selecteren. Hierdoor vindt de Google Lens meer resultaten dan Tineye, maar is het soms wel moeilijker om exact dezelfde afbeeldingen te vinden. Een voordeel van Google Lens is dat het ook kan worden geïntegreerd in de Chrome-browser.

Kosten voor de gebruiker: Gratis beschikbaar.

Pimeyes - EMEARobotics

Functionaliteit: Reverse image search met geavanceerde gezichtsherkenningstechnologie. Pimeyes is vooral geschikt voor het identificeren van afbeeldingen van personen. De software vergelijkt gezichten met verschillende beeld databanken. Onder andere factcheckers van Knack gebruikten de tool voor het identificeren van personen.

Kosten voor de gebruiker: €35,99/maand voor 25 zoekopdrachten per dag, €358,99/maand voor de deep search-functie, waarmee volgens Pimeyes grondiger wordt gezocht.

2.8.5 Vier ethische vragen over AI-tools

Vrijwel alle hierboven besproken programma's kunnen zeer bruikbaar zijn voor journalisten. De vraag in hoeverre het gebruik ervan ook verantwoord is, is moeilijker te beantwoorden. Vooral bij de beeldgeneratiemodellen blijven verschillende ethische vragen onbeantwoord door een gebrek aan transparantie van de bedrijven erachter.

- *Op welke data zijn de modellen getraind?*

Midjourney, Dall-E, Ideogram en Grok onduidelijk over de data waarop hun modellen getraind zijn.

Het bedrijf achter Stable Diffusion gaf wel openheid over de gebruikte trainingsdata, maar bleek te zijn getraind op een beeldbank die grotendeels bestond uit niet-geautoriseerde afbeeldingen. Stability AI werd als een van de eerste AI-bedrijven aangeklaagd voor copyrightscheidingen, onder meer door Getty Images.

Dat plaatst ook vraagtekens bij ‘ethische’ AI-programma’s zoals Tess AI, dat de artiesten wier werk het imiteert weliswaar vergoedt, maar dat fundamenteel gebaseerd is op het Stable Diffusion-model. Alleen het AI-beeldgeneratiemodel van stockfotobedrijf Shutterstock lijkt volledig in orde te zijn met de toestemming van de makers van de trainingsdata van hun modellen. Die makers krijgen ook een deel van de vergoeding.

Het verantwoord verzamelen van trainingsdata is slechts een van de ethische kwesties rond AI-beeldgeneratie. Er zijn nog andere vragen waar niet altijd voldoende duidelijkheid over wordt gegeven:

Hoe gaan de bedrijven achter AI-modellen om met bias (bestrijden)?

Hoe zijn de arbeidsomstandigheden van ‘data labellers’ die modellen trainen en verbeteren door gegevens te annoteren of labelen?

Wat is de klimaatimpact van de programma’s?

Ook Shutterstock heeft hierover nog geen duidelijkheid gegeven.

Wat betreft het gebruik van AI-tools voor **beeldverificatie**, lijkt het antwoord of er verantwoord gebruik van kan worden gemaakt iets eenduidiger. Ook hier is nog weinig informatie beschikbaar over de manier waarop de modellen gemaakt zijn. Maar in tegenstelling tot beeldgeneratieprogramma’s creëren deze tools geen nieuwe beelden en schenden ze dus geen copyright. Bovendien is er een minder arbeidsintensief en energieverslindend trainingsproces nodig om de detectiemodellen te trainen, en blijkt uit onderzoek van Hugging Face en Carnegie Mellon University (Luccioni et al., 2024) dat het gebruik van AI voor beeldverificatie een veel kleinere klimaatimpact heeft dan AI-beeldgeneratie.

Een belangrijke afweging bij het gebruik van **AI-beelddetectietools** is dat geen enkel programma 100% waterdicht is. Dat maakt programma’s als Truepic Vision interessant, die een omgekeerde aanpak hanteren: het certificeren van mensgemaakte beelden. Naast Truepic heeft ook Sony software ontwikkeld die een ‘authenticiteitszegel’ in hun camera’s toevoegt aan de afbeelding, dat zou bewijzen dat de afbeelding niet met AI maar met een optische camera is gemaakt. Echter, zowel Truepic als Camera Authenticity Solution, het systeem van Sony, werken met betaalde abonnementen (prijs op aanvraag). Daarmee dreigt er een financiële drempel te ontstaan voor fotografen die hun menselijkheid willen bewijzen.

Ook AI-software voor **reverse image search** heeft een relatief gezien beperkte klimaatvoetafdruk, en gaat niet gepaard met het risico op copyrightscheiding. Bij deze software kunnen zich echter wel privacyproblemen voordoen. Hoewel zowel Google Lens als TinEye aangeven de geüploade zoekafingen niet langdurig op te slaan, is bekend dat het verdienmodel van Google voor een aanzienlijk deel gebaseerd is op het verkopen van gebruikersdata aan adverteerders. Het track record van Google op het gebied van privacy en mensenrechten maakt van TinEye een aantrekkelijkere keuze voor journalisten.

Pimeyes is dan weer controversieel door de akelig accurate gezichtsherkenningstechnologie, die er volgens critici voor dreigt te zorgen dat niemand meer anoniem over straat kan. Bovendien geeft Pimeyes expliciet aan dat geüploade afbeeldingen ook gebruikt worden voor het verbeteren van hun zoekresultaten. De tool is extreem bruikbaar voor journalisten, maar vraagt dus ook voorzichtigheid en een grondige privacy-afweging.

Hoofdstuk 3

Methode

Methode

Dit hoofdstuk beschrijft de methodologische opzet van het onderzoek naar AI-gegenereerd beeld in de journalistiek. We hanteerden hiervoor een combinatie van kwalitatieve en experimentele onderzoeksmethoden om een diepgaand en breed gedragen inzicht te verkrijgen in de percepties, ethische afwegingen en praktijkervaringen van de diverse belanghebbenden wat betreft AI en beeldgebruik in de nieuwsmedia.

Allereerst wordt de deskresearch toegelicht (3.2). Vervolgens behandelen we de semi-structureerde diepte-interviews afgenomen onder een diverse groep belanghebbenden, waaronder beeldredacteurs, fotojournalisten, experts en mediaombudspersonen (3.3). Hiernaast werd een beperkt experiment uitgevoerd om aanvullende inzichten te verkrijgen over de manier waarop praktijkprofessionals AI-gegenereerd beeld beoordelen (3.4).

3.1 Deskresearch

Een eerste cruciale stap in dit onderzoek was deskresearch, “the process of collating and coding existing information for analysis, without direct contact between researchers and research participants” (Green & Cohen, 2021), dat we uitgebreid beschreven in de literatuurstudie (zie Hoofdstuk 2). Door middel van een analyse van recente academische literatuur hebben we de bestaande kennis in kaart gebracht en een theoretisch kader geboden voor verdere analyse. Aan bod komen achtereenvolgens recente ontwikkelingen in AI en de journalistiek, visuele AI, vertrouwen, (beeld)richtlijnen in de journalistieke media, transparantie en de mogelijkheden en risico's van beeld AI in de journalistiek, en ethische dilemma's met een specifieke focus op de inzet van AI-gegenereerd beeld.

Daarnaast verzamelden we inzichten uit journalistieke bronnen en vakpublicaties zoals Nieman Lab, dat trends en richtlijnen binnen de media-industrie monitort. Door deze bronnen te combineren, brachten we niet alleen de bestaande richtlijnen en hun onderliggende principes in kaart, maar konden we ook lacunes en discussiepunten identificeren, die verdere studie en toetsing in de praktijk vereisen.

3.2 Diepte-interviews

In dit onderzoek hanteerden we een kwalitatieve benadering, waarbij we gebruikmaakten van semigestructureerde diepte-interviews. Deze methode stelde ons in staat om diepgaande inzichten te verkrijgen in de percepties, ervaringen en overwegingen van diverse belanghebbenden met betrekking tot AI-gegenereerd beeld in de journalistiek. De interviews boden de nodige flexibiliteit om dieper in te gaan op de inzichten, interesses en expertisegebieden van de respondenten. De interviews vonden plaats van mei 2024 tot en met januari 2025, zowel fysiek als digitaal via Zoom of MS Teams, met een gemiddelde duur van 60 minuten.

We streefden naar een evenwichtige verdeling van geïnterviewden en besteedden bijzondere aandacht aan de selectie van respondenten, waarbij elk individu over de nodige expertise en kennis beschikt om relevante perspectieven te bieden (tabel 1). Om vertrouwelijkheid te waarborgen, werden de respondenten geanonimiseerd en kregen zij identificatiecodes toegewezen die gekoppeld zijn aan hun functies.

Tabel 1 - Geïnterviewde experts en praktijkmensen

Experts: totaal 19		
14 man, 5 vrouw		
Praktijkmensen: totaal 40		
28 man, 12 vrouw 34 nationaal, 6 regionaal		
Beeldredacteurs	12	8 man, 4 vrouw
Fotografen	12	8 man, 4 vrouw
Grafisch vormgevers	11	9 man, 2 vrouw
Ombudspersonen	5	4 man, 1 vrouw

De uiteindelijke selectie bestaat uit 59 geïnterviewden, die we onderverdeelden in twee hoofdcategorieën: experts en praktijkmensen:

- 19 experts: Deze groep bestaat uit academici, vertegenwoordigers van fotograferingsorganisaties, media- en AI-experts (uit binnen- en buitenland), die beroepsmatig reflecteren op de ontwikkelingen in fotografie en kunstmatige intelligentie. Daarnaast worden juridische experts uit de praktijk geïnterviewd om de juridische implicaties van AI-gegenereerd beeld te verkennen (bijlage 1).

Daarnaast ondervroegen we praktijkmensen, bestaande uit beeldmakers en ombudspersonen (bijlage 2):

- 12 beeldredacteurs: Als direct verantwoordelijken voor de beeldselectie in de journalistiek spelen beeldredacteurs een sleutelrol in het waarborgen van ethische en visuele kwaliteitsnormen binnen nieuwsorganisaties.
- 12 fotografen: Deze groep omvat professionele beeldmakers die, al dan niet met de nieuwste technieken, snel en verantwoord journalistiek beeld produceren. Hun ervaringen en perspectieven op AI-geassisteerde en gegenereerde beelden bieden essentiële inzichten in de praktijk.
- 11 grafisch vormgevers: maken illustraties, ontwerpen en bewaken de huisstijlen van media.
- 5 mediaombudspersonen: Deze respondenten hebben expertise in publieksklachten en ethische vraagstukken binnen de journalistieke praktijk en reflecteren beroepsmatig op innovaties binnen het medialandschap en de bredere samenleving. De mediaombudspersonen zijn geen beeldmakers, maar discussiëren wel in de praktijk met lezers/kijkers en makers over de beeldkwaliteit van de media waar ze voor werken.

In alle groepen waren ruim meer mannen dan vrouwen. Een resultaat dat overeenkomt met de verhouding van leden van de beroepsvereniging NVJ/NVF (ruim 80% man) (Prüst 2025).

Door deze brede en representatieve selectie van belanghebbenden brachten we niet alleen de journalistieke praktijk in kaart, maar ook de juridische, ethische en maatschappelijke implicaties van AI-gegenereerd beeld binnen de media. Publieksonderzoek van buiten het kader van dit project.

Binnen dit onderzoek speelden zoals gezegd wetenschappelijke interviews een centrale rol als kwalitatieve methode, vanwege hun vermogen om diepgaande inzichten te genereren, nuances bloot te leggen en een breder contextueel begrip te ontwikkelen. In tegenstelling tot gestandaardiseerde vragenlijsten biedt deze interactieve benadering de flexibiliteit om in te spelen op de specifieke expertise en ervaringen van respondenten. Dit opent de mogelijkheid voor onverwachte ontdekkingen en verhelderende details die anders onopgemerkt zouden blijven.

Fasen: De interviews verliepen in twee opeenvolgende stappen. In de eerste fase werden gesprekken gevoerd met een zorgvuldig geselecteerd expertpanel, bestaande uit professionals met diepgaande kennis van en ervaring met AI en beeldgebruik in de journalistiek. Deze gesprekken waren gericht op het identificeren van de belangrijkste thema's, uitdagingen en trends binnen dit domein. De inzichten uit deze fase vormden de basis voor de gestructureerde vragen en probleemcases die worden gebruikt in de tweede gespreksronde.

De tweede fase omvatte interviews met fotojournalisten, grafisch vormgevers ('makers'), beeldredacteurs en media-ombudspersonen ('beoordelaars'). In deze gesprekken werd dieper ingegaan op de ethische en praktische implicaties van AI-gegenereerde beelden binnen journalistieke producties. De combinatie van gestructureerde vragen en een probleemcase zorgde ervoor dat de interviews niet alleen praktijkgerichte inzichten opleverden, maar ook bijdragen aan een bredere theoretische reflectie.

De interviewgids bestond uit drie secties:

1. De eerste sectie focuste op algemene demografische informatie van de respondenten.
2. De tweede sectie onderzocht de kansen en gevaren van het gebruik van generatieve beeld-AI-tools.
3. De derde sectie behandelde ethische vraagstukken met betrekking tot het gebruik van generatieve beeld-AI en de ontwikkeling van richtlijnen.

Een onderdeel van deze interviews was het voorleggen van een probleemcase, waarin saillante kwesties en ethische dilemma's met betrekking tot AI-gegenereerd beeld centraal staan. Deze casus is gebaseerd op de deskresearch en eerdere gesprekken met media-partners en dient als referentiekader om de professionele afwegingen en reacties van de respondenten te analyseren.

Analyse. Elk interview, fysiek of digitaal, werd opgenomen. Om een gestructureerde en reproduceerbare analyse te garanderen, werden alle interviews daarna getranscribeerd met behulp van GoodTape of AmberScript, beide AI-tools die audiobestanden automatisch omzetten in geschreven tekst. De transcripten werden vervolgens grondig gecontroleerd, gecodeerd en geanalyseerd met Atlas.TI, waarbij kernideeën werden geïdentificeerd en gelabeld op basis van categorieën en concepten. Deze methode maakt het mogelijk om patronen en thema's binnen de data te identificeren. Om de onderzoeksvragen te beantwoorden, ontwikkelden we een coderingsschema op basis van opkomende patronen die werden ontdekt tijdens het lezen van de interviewtranscripties (Strauss & Corbin, 1990).

De methodologische nauwkeurigheid van dit onderzoek werd gewaarborgd door het systematisch documenteren van de onderzoeksprocedures. De eerste gespreksronde werd vastgelegd in een interviewprotocol, waarin de richtlijnen en werkwijze transparant werden gedocumenteerd. Ook de gestructureerde vragen voor de tweede fase werden zorgvuldig geregistreerd. Deze aanpak draagt bij aan de coherentie en validiteit van het onderzoek en verhoogt de generaliseerbaarheid van de bevindingen.

De expertinterviews werden afgerond binnen de eerste drie maanden van het onderzoek, waarna de tweede fase volgde met een doorlooptijd van vier maanden. De analyse en rapportage zijn in de daaropvolgende twee maanden uitgevoerd, inclusief validatiemomenten met de betrokken mediapartners. Daarnaast werden intervisiemomenten met de mediapartners ingebouwd, waarin de voortgang van het onderzoeksproject werd geëvalueerd en waar nodig methodologische bijsturing plaatsvond. Deze intervisies werden gepland in de voorbereidingsfase, halverwege de interviewperiode en bij de afronding van de analyse.

Met deze systematische en transparante werkwijze wordt niet alleen de wetenschappelijke betrouwbaarheid van het onderzoek gewaarborgd, maar ook een praktijkgericht kader gecreëerd voor een diepgaande en relevante analyse van de inzet van AI-gegenereerd beeld in de journalistiek.

3.3 Experiment

Naast de deskresearch en kwalitatieve interviewfase werd een experimentele component geïntegreerd als aanvullende onderzoeksmethode. Hoewel dit experiment een kleinere rol binnen het totale onderzoeksontwerp vervulde, bood het ons waardevolle ondersteuning bij de interpretatie van onze andere onderzoeksinzichten. De resultaten zijn niet significant in statistische zin, maar dragen bij aan een beter begrip van de percepties en afwegingen van praktijkprofessionals ten aanzien van AI-gegenereerd beeld.

Ongeveer de helft van de betrokken praktijkprofessionals nam deel aan het experiment. Zij kregen twaalf beelden voorgelegd en moesten beoordelen of deze foto's waren of AI-gegenereerde beelden. Ook werd hen gevraagd hardop te denken over hoe ze tot hun oordeel kwamen (think aloud-protocol; Jaspers et al 2004). Door deze benadering werd op een gestructureerde manier onderzocht hoe zij reageren op verschillende vormen van AI-gegenereerde visuele content binnen een journalistieke context. De bevindingen uit dit experiment dienen ter triangulatie met de kwalitatieve interviews en dragen zo bij aan een rijkere duiding van de onderzoeksresultaten.

Hoofdstuk 4

Resultaten

Resultaten

Deze studie onderzoekt de impact van AI op beeldproductie binnen journalistieke media en richt zich op vijf kernvragen:

1. Welke mogelijkheden biedt AI voor beeldproductie?
2. Welke risico's brengt dit met zich mee?
3. Welke richtlijnen hanteren redacties bij het omgaan met AI-beelden?
4. Welke expertise is nodig voor een verantwoorde selectie en productie van AI-generereerd beeld?
5. Welke vormen van transparantie zijn wenselijk om de herkomst en betrouwbaarheid van visueel journalistiek materiaal te verduidelijken?

Door middel van deskresearch en interviews met experts en journalistieke professionals en een kritische reflectie op de ethische implicaties van AI-beeldproductie, biedt dit onderzoek een overzicht van de stand van zaken en de uitdagingen waar nieuwsredacties voor staan. In deze resultatensectie verschaffen we inzicht in hoe journalistieke organisaties AI-beeld inzetten en welke stappen nodig zijn om de transparantie en integriteit van journalistiek beeldmateriaal in het digitale tijdperk te waarborgen.

4.1 Mogelijkheden van journalistieke beeld-AI

Vraag 1: Welke mogelijkheden biedt AI aan beeldproductie voor journalistieke media?

Eerst worden enkele geslaagde voorbeelden beschreven. Uit de interviews blijkt dat AI voornamelijk wordt gebruikt voor efficiëntie (4.1.2), illustratieve toepassingen (4.1.3) en het creëren van beelden die anders moeilijk of onmogelijk te realiseren zijn (bijvoorbeeld vanwege ethische, juridische of veiligheidsbeperkingen) (4.1.4).

4.1.1 'Geslaagde' voorbeelden

Een van de vragen tijdens de interviews was: "Kun je een voorbeeld geven van geslaagde toepassing van beeld-AI in de journalistiek?" Dat leverde opvallend weinig voorbeelden op, zowel bij de ondervraagde experts als de professionals die actief zijn in beeldproductie en -beoordeling. Aangedragen voorbeelden zijn voornamelijk van buitenlandse casussen, maar er waren ook enkele spraakmakende van eigen bodem:

Een expert noemt een **graphic novel over de inval in Oekraïne** – een digitaal, grafisch stripverhaal – dat door de BBC met behulp van AI is gecreëerd. Volgens de respondent experimenteerde de Britse omroep met een nieuwe aanpak: de dagelijkse schriftelijke verslaggeving van de Russische inval in Oekraïne werd ingevoerd in een AI-beeldgenerator. Zo ontstond een visuele liveblog, waarbij de oorlog in realtime werd vertaald naar indrukwekkende, door AI gegenereerde beelden. Elke dag werd de voortgang van het conflict dus vastgelegd in een visuele update, gebaseerd op de teksten van BBC-verslagvers.

Een expert spreekt over de ‘Michael Christopher Brown-case’. Brown, zelf van origine fotojournalist, gebruikte de generatieve AI-tool Midjourney om een fotorealistische ‘post-photography reportage’ te maken over de grootschalige, dertig jaar durende Cubaanse migratie van Havana naar de Verenigde Staten, waarvan geen beeld beschikbaar was. Meer dan 25 jaar lang hield Brown een lijst bij van onderwerpen die hij wilde documenteren maar niet kon documenteren, voornamelijk vanwege de toegang of omdat toegang onmogelijk was.

Over zijn project schrijft Brown:

“Publicaties zoals The New York Times en The New Yorker gebruiken al tientallen jaren reportage-illustraties van verschillende soorten naast verslaggeving, om ervaringen, ideeën en mensen op verschillende manieren te belichten en ons idee van een onderwerp uit te breiden. Er zijn dingen die kunstmatige intelligentie kan toevoegen aan de beeldindustrie, die zeer nuttig kunnen zijn, en dat is de plek waar we moeten kijken.”

Een grafisch ontwerper haalt een productie van diens eigen medium, het Brabants Dagblad, dat geregeld AI gebruikt bij de totstandkoming van uitlegvideo's. Een voorbeeld is een explainervideo over de Delta Rijn Corridor, waar AI is ingezet voor grafische elementen, zoals een doorsnede van de pijpleiding, maar ook voor fotorealistische verbeeldingen van de omgeving, die dienen als achtergrond. Bomen, een hijskraan, graafwerkzaamheden: volgens de respondent allemaal gegenereerd met Adobe Firefly. De respondent onderbouwt diens keuzes als volgt:

“Ik probeer zelf B-roll te schieten. Dat is extra footage dat in je video zit, maar dat lukt niet altijd. We kunnen niet altijd op locatie gaan, dus zoeken we dan stockvideo's op. We zijn partners met het Belgische ‘Het Laatste Nieuws’, die hebben een hele database met stockvideo's. Bij Adobe hebben ze natuurlijk ook stockvideo's, maar ja, dat geeft ook niet altijd – precies wat je wil en past soms net niet. Ja, en dan proberen we dit gewoon een keer. Dit was in de tijd gemaakt dat het een beetje in opkomst was, en dit is het resultaat.”

Een praktijkrespondent haalt een voorbeeld van diens eigen medium (AD) aan: een met AI-gegenereerde illustratie boven een AD-column met de titel ‘Ik vond mezelf zwevend boven de binnenstad van Doetinchem’. De AI-illustratie toont inderdaad een zwevende man boven de binnenstad van Doetinchem. Een respondent meent dat een redacteur de column – geschreven door een freelancer – van beeld voorziet. De respondent zegt het volgende over het productieproces:

“Het is natuurlijk maar net of je dus letterlijk die hele tekst in een AI-tool gooit, of alleen een paar zinnen. Geen idee in dit geval. Maar op die manier wordt er dus al gebruik gemaakt van AI, op zich op een illustratieve manier. Weet je, hier is geen misverstand over mogelijk dat het een foto is of een illustratie, maar omdat er zoveel columnisten zijn, zoveel regio's, ga je niet iedere column door een illustrator op deze manier laten illustreren. Dit [beeld-AI bij columns, red.] is iets dat ik veel vaker zie gaan gebeuren.”

Het meest aangehaalde voorbeeld is dat van **AI-nieuwslezer Nina van Omroep Brabant**. Nina van den Broek is een ‘echte’, menselijke verslaggever in dienst van Omroep Brabant wiens stem en uiterlijk met AI zijn gekloond. De door journalisten geschreven nieuwsberichten worden dus door de AI-Nina gepresenteerd, op een natuurgetrouwe manier in stem en houding. Na uitgebreide wederzijdse afspraken en verantwoording over het hoe en waarom van dit innovatieve experiment, startte Omroep Brabant in mei 2024 met het uitzenden van de AI-versie van Nina. Dit experiment wordt door respondenten besproken, waarbij de bezorgdheid overheerst. Een expert:

“Je geeft een boodschap aan het publiek. Vertrouw ons maar. Ook al ziet u hier een machine en niet een mens. En ik denk dat je het andersom moet doen. Want je moet zeggen, ik als mens vertel dat u mij moet vertrouwen en ik heb daar wel allerlei dingen bij gebruikt. Ik vind het te makkelijk om te zeggen: het is een schande wat Omroep Brabant doet. Ik denk voor de langere termijn qua geloofwaardigheid van hun merk, doen ze er verstandig aan om aan het eind van het proces de mens te laten werken. En waarom? Omdat het publiek uiteindelijk ook instinctief een weg moet zoeken naar wat het nog wel en niet kan geloven.”

De hoofdredacteur van Omroep Brabant verweert zich tegen afwijzingen van dit experiment, met daarin AI-nieuwslezer Nina, als volgt:

“Onze uitgesproken teksten zijn wel door journalisten geschreven. Wat is het verschil met de NOS-nieuwsapp, waarin de berichten van de NOS worden voor-gelezen? Dat is ook een door AI-gegenereerde stem, die de berichten van de NOS voorleest. NOS doet toch exact hetzelfde wat wij doen, alleen we hebben er een plaatje bij, dat jou vertelt wat er aan de hand is? Het zit blijkbaar in het beeld.”

Het experiment is intussen [beëindigd](#): “te bewerkelijk, te duur en te verwarrend voor de kijkers”. De boodschap die de omroep aan het begin van de filmpjes in beeld bracht (‘De presentator van deze video is met AI gemaakt. Al het nieuws en de beelden zijn natuurlijk echt.’) was blijkbaar niet duidelijk genoeg, want kijkers vroegen zich geregeld af naar welke Nina ze zaten te kijken: de echte of de AI-kloon. Op Facebook regent het negatieve commentaren: van “walgelijk, dat AI” tot “bezopen zooi” en ten slotte “weg gaat de mensheid we worden vervangen door robots. Vind het levensgevaarlijk.”

Opvallend dat bijna niemand De Geïllustreerde Kunstmatige Intelligentie: AI-Journalistiek in Actie noemt. Dit is een site met volledig geautomatiseerde nieuwsgaring en -publicatie, zonder menselijk ingrijpen van begin tot eind. Op [dgki.nl](#) staat te lezen dat deze site het podium is waar Laio (de AI powered redacteur van Innovation Origins) haar kunnen toont:

“Laio is meer dan een verzameling algoritmes; het is een symbiose van technologie en journalistiek inzicht. Hierin neemt Laio het voortouw in het vinden, pitchen en schrijven van verhalen die innovatie en technologische vooruitgang onder de aandacht brengen. Met haar taalcompetenties kan Laio nieuws interpreteren en de meest substantiële inhoud voor publicatie identificeren.

Hoewel AI de potentie heeft om het journalistieke landschap ingrijpend te veranderen, vervangt het de menselijke journalist niet. Machines overstijgen in precisie en analysevermogen, maar het is de mens die verhalen leven inblaast. Maar zonder mensen is al veel mogelijk. Dat te laten zien is het doel van dgki.nl. [...]

Het is belangrijk te benadrukken dat ‘De Geïllustreerde Kunstmatige Intelligentie’ functioneert als een demonstratieplatform, niet als een gevestigde nieuwsautoriteit. Het reflecteert het potentieel en legt de uitdagingen van AI in de journalistiek bloot, inclusief inherente vooroordelen en foutmarges die momenteel nog bestaan.”



Estafettetakingen over het hele land gestart om bezuinigingen in hoger onderwijs aan te vechten (AI-gegenereerde illustratie bij nieuws van dinsdag 11 maart 2025 op dgki.nl)

De beperkte oogst aan goede voorbeelden wijst erop dat er binnen het veld nog geen breed gedeeld referentiekader bestaat voor succesvolle AI-toepassingen in journalistieke beeldproductie. Verder valt op dat vooral de experts weinig voorbeelden gaven van AI-tools waarmee ze zelf ervaring hadden. Dit staat in contrast met beeldmakers en beeldredacteurs, die, zoals verwacht, een breder scala aan applicaties kennen en gebruiken.

AI-tools worden door de respondenten vooral ingezet voor beeldbewerking, zoals het aanpassen van kleuren, scherpte of contrast. Opvallend is dat respondenten alleen Adobe Photoshop – met daarin de onvermijdelijke AI-functionaliteiten – gebruiken voor beeldbewerking en selectietaken.

Wat betreft andere tools, blijkt dat er een onderscheid wordt gemaakt tussen bewerkingsmogelijkheden en programma's die volledig nieuwe beelden genereren. Generatieve AI wordt met meer terughoudendheid benaderd, vooral in journalistieke contexten. Veelgenoemde en vaakgebruikte 'generatietools' zijn Dall-E, MidJourney en Adobe Firefly, onderdeel van Photoshop.

In generatieve AI wordt, zeker bij landelijke media, nog met voorzichtigheid geëxperimenteerd, aldus de respondenten. Bij regionale media, waaronder bij Omroep Brabant en regionale AD-titels, gebeurt dit vaker.

4.1.2 Routinematige efficiëntie

Hoewel onze respondenten slechts enkele goede voorbeelden kunnen opnoemen van het gebruik van generatieve AI bij beeldcreatie, worden er wel specifieke voordelen erkend voor de journalistieke praktijk. Een van de meest gewaardeerde toepassingen is de routinematige efficiëntie die AI biedt. De voornaamste voordelen hiervan zijn tijdsbesparing, workflow-optimalisatie en de mogelijkheid om complexe of moeilijk te verkrijgen beelden te genereren. Een beeldredacteur zegt daarover: "Met bewerkingen die dan wel kwalitatief genoeg zijn, kunnen we best wel voor een ontlasting zorgen bij de beeldbewerkers, waardoor zij ook minder basale handelingen moeten uitvoeren en ook weer creatief te werk kunnen gaan."

Praktijkmensen benadrukken dat AI een waardevol hulpmiddel kan zijn bij het uitvoeren van repetitieve taken, zoals het snel genereren van schetsen, eenvoudige visuals of concepten. AI-tools worden ook vaak gebruikt voor beeldbewerkingstaken zoals het aanpassen van kleuren, contrast en scherpte. Selectietools op basis van AI helpen dan weer om sneller bepaalde beeldelementen te isoleren of aan te passen zonder tijdrovend handmatig werk:

"In een deel van de workflow zijn er AI-tools die heel handig zijn, vooral selectietools. Wil je de lucht iets donkerder maken, dan kan dat met twee klikken. Dit versnelt werk dat anders uren zou duren."

Dit laat professionals meer tijd over voor creativiteit en inhoudelijk werk. Zoals een respondent opmerkte:

“Als ik aan mijn eigen werk denk, is het echt meer een hulpmiddel om het leven gemakkelijker te maken en je workflow te versnellen, zodat je je kunt richten op meer inhoudelijke taken.”

Deze praktische inzet van AI ondersteunt de journalistieke workflow zonder de kernwaarden van het vak aan te tasten.

Een expert ziet, in het kader van routinematige efficiëntie, de voordelen van een gekloonde nieuwslezer (de AI-Nina van Omroep Brabant): dat scheelt veel tijd die nu aan verslaggeving besteed kan worden. Toch bleek de technische uitvoering van dit experiment tijdrovend en duur.

“Een verfijning van de techniek moet het goedkoper maken om de AI-presentator in te zetten. In de ideale situatie worden nieuwsteksten die toch al zijn geschreven met één druk op de knop omgezet in een filmpje waarin een AI-presentator de woorden uitsprekt.”

Daar wordt achter de schermen nu aan gewerkt, maar tot die tijd staat het Nina-experiment stil.

Opmerkelijk: de voordelen die AI kan bieden om routinematige klussen efficiënt en snel uit te voeren, worden breed uitgesproken. Toch blijkt in deze praktijk dat de winst soms nog beperkt is.

4.1.3 Mogelijkheden voor illustratie

Generatieve AI wordt steeds vaker ingezet voor de productie van illustratieve en abstracte beelden, met name in situaties waarin geen geschikt beeldmateriaal beschikbaar is. Dit geldt bijvoorbeeld voor conceptuele onderwerpen zoals inflatie, zelfdoding of klimaatverandering, waarbij het visueel vertalen van complexe fenomenen een uitdaging vormt. Dan kan AI-beeldgeneratie een waardevol instrument zijn om abstracte concepten op een toegankelijke manier te illustreren.

- Als er geen concreet beeldmateriaal beschikbaar is: “Soms als ik een afbeelding zoek en niet weet wat ik moet gebruiken, zoals inflatie, dan gebruik ik AI-beeldgenerators om ideeën op te doen of een bruikbare afbeelding te genereren.”
- Als een kostenbesparend alternatief voor stockfoto's of handgemaakte illustraties. In plaats van het aankopen of laten ontwerpen van visueel materiaal, kunnen redacties met AI-tools snel en efficiënt beelden genereren die aansluiten bij de inhoud van hun berichtgeving.
- Als complexe fenomenen verbeeld moeten worden. Ook binnen audiovisuele producties, zoals documentaires, wordt AI steeds vaker toegepast als een krachtig visueel hulpmiddel. Een expert benadrukt de meerwaarde voor het verbeelden van complexe fenomenen: “AI is nuttig voor visualisaties in documentaires, zoals het effect van klimaatverandering, wat sterker werkt dan een grafiek.”

4.1.4 Kansen voor ‘onmogelijke’ beelden

AI biedt volgens de respondenten ook de mogelijkheid om visuele content te genereren in situaties waarin traditionele beeldproductie niet haalbaar is vanwege ethische, juridische of veiligheidsbeperkingen.

- Als privacy fotograferen onmogelijk maakt. Dit speelt bv. een belangrijke rol bij onderwerpen waarbij privacywetgeving, zoals de Algemene Verordening Gegevensbescherming (AVG), het gebruik van bestaand beeldmateriaal beperkt. Denk daarbij aan beelden van minderjarigen. In dergelijke gevallen kan AI worden ingezet om personen onherkenbaar in beeld te brengen zonder inbreuk te maken op hun privacy.
- Als het te gevaarlijk is. Een expert uit de sector benadrukt het potentieel van AI in deze context zonder dat journalisten of betrokkenen aan risico's blootgesteld hoeven worden: "AI biedt kansen voor beelden die anders moeilijk te maken zijn, zoals bij gevaarlijke locaties, bij conflictgebieden of rampensituaties."
- Bij historische reconstructies en toekomstbeschrijving. Op basis van archiefmateriaal of beschrijvingen, kunnen journalistieke producties visuele representaties creëren van gebeurtenissen of locaties die niet (langer) gefotografeerd kunnen worden. Daarnaast kan AI worden ingezet om de toekomst te verbeelden, bijvoorbeeld door mogelijke scenario's van klimaatverandering, stedelijke ontwikkeling of technologische evoluties visueel vorm te geven.

4.1.5 Kansen voor satire

Beeld-AI rukt op in de artistieke wereld en commerciële wereld en in de media waarin satire en parodie een rol spelen, bijvoorbeeld in de humorpagina's van de kranten of in de programma's van Lubach en Even tot hier. Respondenten herkennen dat en maken duidelijk dat dat wat hen betreft ook geen journalistiek is. Een expert:

"Je wil bijvoorbeeld een parodie maken. Je wil een spotprent maken. Dan is alles geoorloofd. Dan mag je juist manipuleren en met de werkelijkheid aan de haal gaan. Maar dan moet het natuurlijk wel duidelijk zijn dat je dus juist probeert om een opinie te verkondigen. Maar dat is niet wat de doorsnee journalist doet. Ik zie even een journalist als iemand die gewoon verslag doet van wat er gaande is in de wereld."

Een andere expert:

"Ik zie het veel meer in de artistieke en commerciële wereld, daar zie ik het zeker gebeuren. Maar in de journalistiek blijft iedereen er gewoon lekker nog vandaan, behalve dan bij het oprukken van de illustratie. Maar die heeft ook niet die werkelijkheidssuggestie."

Ook hij ziet geen kansen binnen de journalistiek, maar wel voor columns en satire:

“We hebben het over het gebruik van manipulatieve ingrepen. En ik denk dat journalisten het daar vrij snel eens zullen zijn dat ze dat gewoon nooit, never nooit zullen toestaan en mogen toestaan. [...] Maar een columnist of iemand die satire maakt. Dan zou je in principe alles [met beeld-AI-middelen, auteurs] mogen doen wat tot je beschikking staat, als je maar duidelijk maakt: ‘Dit is satire.’”

4.2 Beeldbewerking versus beeldmanipulatie

De respondenten in deze studie maakten een duidelijk onderscheid tussen ‘beeldbewerking’ en ‘beeldmanipulatie’, een kwestie die, naar hun zeggen, in de context van AI-generende beelden steeds urgenter wordt. **Beeldbewerking** omvat aanpassingen die de zichtbaarheid en esthetische kwaliteit van een afbeelding verbeteren, zonder de feitelijke inhoud of betekenis ervan te veranderen. Dit kan bijvoorbeeld het aanpassen van kleuren, contrast en scherpte omvatten, waardoor de aandacht van de kijker beter wordt geleid zonder de realiteit te vervormen.

Beeldmanipulatie daarentegen impliceert het actief wijzigen van de realiteit, wat de journalistieke neutraliteit en geloofwaardigheid in gevaar brengt. Dit gebeurt wanneer AI wordt ingezet om elementen uit een beeld te verwijderen, toe te voegen of te verplaatsen, wat kan leiden tot een vervormde of misleidende representatie van de werkelijkheid. Dergelijke praktijken ondermijnen de waarheidsgetrouwheid van journalistieke content en vergroten het risico op desinformatie. Generatieve AI maakt deze kwestie nog complexer, aangezien AI niet alleen bestaande beelden kan aanpassen, maar ook volledig synthetische beelden kan creëren. Hierdoor ontstaat een grijs gebied waarin het onderscheid tussen bewerking en manipulatie vervaagt.

Een fotograaf beschrijft dit verschil saillant:

“Tegenwoordig kan je dus een onderwerp selecteren of een achtergrond selecteren. En dat doe ik eigenlijk vrijwel standaard om het één naar voren te halen of het ander wat naar achter, of daar wat verschil in te krijgen of de kleur wat te veranderen. Omdat ik altijd met daglicht werk [...] Ik denk nog steeds dat dat mag, want je mag dingen zwart wit maken. Je mag de kleur iets roder maken. Je mag alleen geen pixels weghalen en niet echt dingen weghalen door te vignetteren (red. Het donkerder maken van de hoeken/randenvan een foto).”

Het hanteren van een strikte scheiding tussen beeldbewerking en manipulatie is daarom volgens de respondenten bevorderlijk voor de geloofwaardigheid en ethische integriteit van de journalistiek.

4.3 Ruimte voor experiment op regionale redacties

Opvallend is dat regionale journalistieke media meer ruimte lijken te hebben voor experimenten met beeld-AI dan hun nationale tegenhangers. Omroep Brabant was bijvoorbeeld de eerste Nederlandse nieuwsorganisatie die een echte presentator via AI kloonde voor videoproductie (zie 4.1.1). Deze innovatieve toepassing had als doel snelle, eenvoudige en kostenefficiënte videocontent te genereren voor de website en app van de omroep. Door de stem en het beeld van een journalist met AI te repliceren, kon de omroep efficiënter korte nieuwsfragmenten produceren zonder telkens een fysieke opname te vereisen.

Een belangrijk motief achter dit experiment was volgens de omroep het aanspreken van jongere doelgroepen en mensen met een lagere leesvaardigheid, die video verkiezen boven geschreven nieuws. Door AI-gegenereerde presentaties in te zetten, probeerde Omroep Brabant nieuws toegankelijker en visueel aantrekkelijker te maken. Om de geloofwaardigheid van deze aanpak te waarborgen, bleef de redactionele controle volledig in handen van menselijke journalisten. Zij bepaalden de inhoud en productie van het nieuws, terwijl de AI-tool slechts werd ingezet als ondersteunend middel. Daarnaast zorgde de omroep voor transparantie over het gebruik van AI, wat bijdroeg aan het vertrouwen van het publiek in de betrouwbaarheid van de nieuwsvoorziening.

Een tweede voorbeeld van regionale innovatie komt van het Brabants Dagblad, dat geregeld AI gebruikt bij de totstandkoming van uitlegvideo's. Bijvoorbeeld de explainervideo over de Delta Rijn Corridor, waar AI is ingezet voor grafische elementen, zoals een doorsnede van de pijpleiding, maar ook voor fotorealistische verbeeldingen van de omgeving, die dienen als achtergrond, gegenereerd met Adobe Firefly.

Een derde voorbeeld vormt het Parool, met hun experimentele AI-zomerserie 'Op en top Amsterdams'. Twee beeldredacteurs vroegen wekelijks aan AI-applicatie MidJourney om Amsterdam te genereren op een kunstzinnige manier, om te kijken tot wat AI nou eigenlijk in staat is: niet tot veel, volgens de makers. 'Subtiliteit kent het machientje niet. Bij 'file' staan de auto's op elkaar.'

Een verklaring zou kunnen zijn dat er meer durf zit bij (enkele) regionale omroepen omdat ze dichter bij hun publiek staan, waardoor ze vertrouwen hebben dat hun experimenten en hun uitleg daarvan goed ontvangen wordt. Ook zijn ze doorgaans goed bereikbaar als het publiek hier vragen of meningen over heeft. Een mogelijke andere verklaring voor deze grotere openheid is dat regionale omroepen vaak met kleinere redacties en beperktere financiële middelen werken, waardoor AI als een kostenbesparend alternatief aantrekkelijk kan zijn (al bleek de AI-nieuwslezer geen financieel voordeel op te leveren, en is het experiment daarom gestaakt). Daarnaast bedienen regionale media een specifiek, lokaal publiek dat mogelijk minder kritisch staat tegenover AI-gebaseerde innovaties dan het bredere, nationale publiek. Waar nationale nieuwsorganisaties bovendien onder sterkere publieke en politieke controle staan en aan strengere journalistieke normen moeten voldoen, ervaren regionale omroepen doorgaans minder institutionele druk. Dit kan hen meer flexibiliteit geven om nieuwe technologieën te testen en toe te passen zonder directe reputatieschade of brede maatschappelijke discussie.

Opvallend is dat regionale journalistieke media meer ruimte lijken te hebben voor experimenten met beeld-AI dan hun nationale tegenhangers. Omroep Brabant was bijvoorbeeld de eerste Nederlandse nieuwsorganisatie die een echte presentator via AI klonde voor videoproductie (zie 4.1.1). Deze innovatieve toepassing had als doel snelle, eenvoudige en kostenefficiënte videocontent te genereren voor de website en app van de omroep. Door de stem en het beeld van een journalist met AI te repliceren, kon de omroep efficiënter korte nieuwsfragmenten produceren zonder telkens een fysieke opname te vereisen.

Een belangrijk motief achter dit experiment was volgens de omroep het aanspreken van jongere doelgroepen en mensen met een lagere leesvaardigheid, die video verkiezen boven geschreven nieuws. Door AI-gegenereerde presentaties in te zetten, probeerde Omroep Brabant nieuws toegankelijker en visueel aantrekkelijker te maken. Om de geloofwaardigheid van deze aanpak te waarborgen, bleef de redactionele controle volledig in handen van menselijke journalisten. Zij bepaalden de inhoud en productie van het nieuws, terwijl de AI-tool slechts werd ingezet als ondersteunend middel. Daarnaast zorgde de omroep voor transparantie over het gebruik van AI, wat bijdroeg aan het vertrouwen van het publiek in de betrouwbaarheid van de nieuwsvoorziening.

Een tweede voorbeeld van regionale innovatie komt van het Brabants Dagblad, dat geregeld AI gebruikt bij de totstandkoming van uitlegvideo's. Bijvoorbeeld de explainervideo over de Delta Rijn Corridor, waar AI is ingezet voor grafische elementen, zoals een doorsnede van de pijpleiding, maar ook voor fotorealistische verbeeldingen van de omgeving, die dienen als achtergrond, gegenereerd met Adobe Firefly.

Een derde voorbeeld vormt het Parool, met hun experimentele AI-zomerserie 'Op en top Amsterdams'. Twee beeldredacteurs vroegen wekelijks aan AI-applicatie MidJourney om Amsterdam te genereren op een kunstzinnige manier, om te kijken tot wat AI nou eigenlijk in staat is: niet tot veel, volgens de makers. 'Subtiliteit kent het machientje niet. Bij 'file' staan de auto's op elkaar.'

Een verklaring zou kunnen zijn dat er meer durf zit bij (enkele) regionale omroepen omdat ze dichterbij hun publiek staan, waardoor ze vertrouwen hebben dat hun experimenten en hun uitleg daarvan goed ontvangen wordt. Ook zijn ze doorgaans goed bereikbaar als het publiek hier vragen of meningen over heeft. Een mogelijke andere verklaring voor deze grotere openheid is dat regionale omroepen vaak met kleinere redacties en beperktere financiële middelen werken, waardoor AI als een kostenbesparend alternatief aantrekkelijk kan zijn (al bleek de AI-nieuwslezer geen financieel voordeel op te leveren, en is het experiment daarom gestaakt). Daarnaast bedienen regionale media een specifiek, lokaal publiek dat mogelijk minder kritisch staat tegenover AI-gebaseerde innovaties dan het bredere, nationale publiek. Waar nationale nieuwsorganisaties bovendien onder sterkere publieke en politieke controle staan en aan strengere journalistieke normen moeten voldoen, ervaren regionale omroepen doorgaans minder institutionele druk. Dit kan hen meer flexibiliteit geven om nieuwe technologieën te testen en toe te passen zonder directe reputatieschade of brede maatschappelijke discussie.

4.4 Risico's voor beeld-AI

Welke risico's bestaan bij gebruik van AI bij beeldproductie voor journalistieke media?

4.4.1 Ethische bezwaren

De respondenten zien veel risico's en benoemen een aantal ethische bezwaren in het gebruik van beeld-AI.

- *Zonder toestemming of vergoeding visueel materiaal van derden gebruiken (copyright/content rights)*

Een vaak aangestipt ethisch knelpunt van het gebruik van generatieve beeld-AI is de onvermijdelijke schending van copyright of zogeheten content rights van duizenden, zo niet miljoenen beeldmakers. Generatieve AI-modellen trainen zich zonder toestemming op visueel materiaal dat van het internet wordt gescraped. Makers van dat materiaal ontvangen doorgaans geen financiële vergoeding, compensatie of andere credits voor het gebruik van hun werk.

Respondenten hekelen de copyrightschiending unaniem en iedereen is voor eerlijke compensatie, maar niemand oppert een totale boycot zolang dit niet is geregeld. Een fotograaf beschrijft zijn gedachten over de makers-kant:

“Ik vind wel dat je de makers daarvoor moet vergoeden, eigenlijk. Want zonder al die honderden miljoenen, miljarden beelden die AI gebruikt om te leren, ben je nergens. Mensen die er hun geld mee verdienen zien hun beelden terug in zo'n AI-generator. Ik vind wel dat je mensen ervoor moet compenseren.”

Welke ethische redenering de gebruiker van een visuele AI-generator zou moeten maken, is volgens dezelfde fotograaf lastig te bepalen:

“Als je iets maakt met zo'n ding... Dat vind ik heel moeilijk: voor een deel is er eigen creativiteit die je in het ding stopt door je prompt goed te beschrijven, en tegelijkertijd maakt een programma jouw ding. Met behulp van Photoshop kan je ook veel dingen en dan blijft het je eigen copyright. Maar dit vind ik wel moeilijk ja. Ik heb hier niet echt een antwoord op.”

Op organisatieniveau is er ook onduidelijkheid over hoe om te gaan met schending van copyright en content rights. Een expert licht dat (hypothetisch) toe:

“Op het moment dat je als organisatie zegt, daar gaan wij gebruik van maken, dan moet je je dus afvragen: in hoeverre zijn wij bereid bedrijfsrisico's, maar ook meer maatschappelijk-ethische risico's te nemen? Je kunt zeggen, ja, dat copyright is niet mijn probleem. Ik wil gewoon een leuk product maken en als ik dat aan Midjourney of Dall-E vraag, dan is dat zo gefixt. Dus ja, is dat mijn probleem? Nee. Die maatschappelijke verantwoordelijkheid zie ik niet. Ik wil gewoon geld verdienen.”

Een eventueel gebrek aan maatschappelijke verantwoordelijkheid over het gebruik van visuele AI-generatoren wordt, tot dusver, juridisch niet aan banden gelegd. Er is geen wetgeving over copyright- of content rights-schending. Een expert zegt daarover:

“As far as I know in the EU the situation is pretty clear that there is no copyright on automatically generated pictures and there is no way to protect your pictures from being used in the system.”

Toch blijkt dat laatste, in elk geval in theorie, onjuist. Er is zoiets als een opt-out initiatief, waar een handjevol respondenten over spreekt. Een van hen werkte mee aan de ontwikkeling ervan en zegt het volgende:

“Er is nooit toestemming gevraagd voor dat scrapen, en waarschijnlijk denken de meeste mensen ook dat dat niet hoeft. [...] In de wet heet text en datamining TDM. Dat zegt eigenlijk dat dat zelfs voor commerciële doeleinden mag. Je mag materiaal verzamelen, reproduceren, die in zo’n dataset zetten en daar een bot mee laten trainen. Dat is allemaal toegestaan, maar dat wil niet zeggen dat je niet toch zou moeten afdragen aan de makers. [...] Dus daar zit onze focus: er moet uiteindelijk vergoeding worden betaald. Dat is één kant van het verhaal. De andere kant is dat heel veel makers niet weten dat je een ‘rechten voorbehoud’ kunt maken. Dus in principe is het toegestaan op grond van die TDM-uitzondering, maar je kunt wel als maker zeggen ‘van mijn beeld blijf je af. Wij hebben een tool ontwikkeld, die heet Opt Out Now, en daarin kun je makkelijk, snel en gratis zo’n voorbehoud instellen. Dat wil eigenlijk zeggen dat die bots het web afstruinen – als kleine spinnetjes eigenlijk – en op een gegeven moment zo’n rechtenvoorbehoud tegen komen. En die zien: hé, wacht eens, ik kan het beeld van meneer Jansen niet gebruiken – en die zouden dat dan moeten overslaan.”

- ***Visuele laster door produceren van AI gemanipuleerd beeld (en geluid)***

Een ethisch risico dat toeneemt met de ontwikkeling van generatieve AI is het gemak waarmee deepfakes gemaakt kunnen worden: fotorealistische verbeeldingen van echte mensen in een kunstmatige situatie en/of positie. Alhoewel Nederlandse mainstream media geen deepfakes genereren of (doelbewust) verspreiden, heeft de journalistiek wel te maken met de circulatie van nepbeelden. Een expert licht dat toe:

“Er is een verschil tussen mainstream media en sociale media, natuurlijk. Ik ga er vanuit dat grote beeldredacties een soort check doen bij beelden om te verifiëren of een beeld echt is. [...] Er is een onderscheid tussen echte journalistiek en misleidende deepfakes. Er worden natuurlijk ook veel seksueel getinte beelden gemaakt, meestal om mensen individueel op te lichten. [...] Ik denk dat het aan de orde van de dag is dat consumenten en particulieren op sociale media op allerlei manieren worden gemanipuleerd en misleid met beeld.”

Een beeldredacteur waarschuwt voor gemanipuleerde beelden (en teksten) van echte mensen die hen schaden:

“Dat is net als met quotes. De Telegraaf maakt er een andere quote van dan het Leidsch Dagblad. Dat is toch een beetje ingrijpen in iemands persoonlijke levenssfeer. Ik heb iets gezegd, in het openbaar, en dat is gewoon verdraaid. Dat wordt misschien wel veel erger.”

“Mag je iemands stem vervormen? Dan moet je wetgeving maken met een lijst wat je niet en wel mag. Je kan misschien wel in de privacywetgeving opnemen dat je iemand niet iets anders mag laten zeggen dan hij heeft gezegd.”

- ***Potentiële misleiding van publiek door suggestie van integer, fotorealistisch beeld***

Het gebruik van AI bij de productie van journalistieke beelden roept ernstige ethische vragen op, voornamelijk met betrekking tot waarheidsgetrouwheid van nieuwsorganisaties. Een van de grootste zorgen die experts en praktijkmensen benadrukken, is dat AI-gegenereerde beelden, zelfs wanneer ze transparant worden gelabeld, door het publiek alsnog als authentiek kunnen worden geïnterpreteerd. Dit risico wordt vergroot door het feit dat mensen vaak onderschriften en disclaimers niet lezen, en na verloop van tijd vergeten dat een beeld door AI is gegenereerd. Zoals een respondent opmerkt:

“Mensen zijn niet zo genuanceerd in hun geheugen. Als het een tekening was geweest of een schilderij dan zien mensen dat het een afbeelding van de werkelijkheid is, maar bij AI-beelden verdwijnt dat besef.”

Daarnaast bestaat het risico dat het gebruik van AI in beeldproductie leidt tot de eerdergenoemde glijdende schaal, waarbij steeds minder onderscheid wordt gemaakt tussen beeldbewerking en beeldmanipulatie (zie 4.3). Terwijl het aanpassen van contrast en kleur binnen journalistieke normen valt, vormt het verwijderen of toevoegen van beeldinformatie een bedreiging voor de journalistieke integriteit.

- ***Fotorealistische beeld-AI tast het vertrouwen in de journalistiek aan***

Veel respondenten tonen grote bezorgdheid over het bredere maatschappelijke effect van AI-gegenereerde beelden op het vertrouwen in journalistiek. Een expert benadrukt:

“Als je een nieuwsorganisatie bent die zich committeert aan waarheid, dan is het gebruik van een stochastisch model dat niets met waarheid te maken heeft per definitie problematisch.”

Dit illustreert een fundamenteel dilemma: AI kan een krachtig hulpmiddel zijn, maar zodra het wordt ingezet om beelden te genereren die niet direct corresponderen met de realiteit, kan dit de geloofwaardigheid van nieuwsmedia ondermijnen.

Ook het wennen van het publiek aan de mogelijkheid dat AI nieuws kan presenteren (de casus Nina) levert het risico van wantrouwen op, volgens een expert: “Omdat het publiek uiteindelijk ook instinctief een weg moet zoeken naar wat het nog wel en niet kan geloven.”

Zowel experts als praktijkmensen beschrijven het belang van het behouden van journalistieke integriteit en publiek vertrouwen door ethische uitdagingen rond AI aan te pakken.

- ***AI-beeld reproduceert vooroordelen***

De opvatting dat AI-generatoren beelden met vooroordelen produceren is door meerdere respondenten uitgesproken. Dat komt volgens hen omdat AI-modellen zich trainen op materiaal waarin onvermijdelijk ook vooroordelen voorkomen. Een beeldredacteur schetst een voorbeeld:

“Toevallig stond er pas een foto in Sportwereld. Het ging over een groep rechters die beslist over het wederom toelaten van Vitesse tot de betaalde voetbalcompetities. Er was een door AI gegenereerde foto bij gezet van vijf rechters in de rechtszaal, vanaf de rug gezien. Het waren vijf grijze mannenhoofden. Mensen vroegen wat er mis mee was. Ik zei: wij geven de indruk alsof het een echte foto is. Wij geven de lezer de indruk dat vijf grijze mannen beslissen [...]. Waarom geen vijf blonde of bruine vrouwen? Waarom vijf grijze mannen?”

Respondenten spreken over bias, stereotypes en vooroordelen die generatieve AI produceert. Een respondent oppert optimistisch:

“Misschien kan het er zelfs voor zorgen, dat hoe we dingen gaan visualiseren, dat dat diverser wordt. Ik zie dat op dit moment nog niet, omdat er in die systemen toch veel bias zit. Maar misschien zorgt het er ook voor dat mensen met een andere huidskleur vaker in beelden terechtkomen, terwijl dat nu vaak veel te weinig is.”

4.4.2 Bezwaren van ‘kennissilo’s’

Een bezwaar van een andere aard dat vaak werd genoemd was dat van het ontstaan van kennissilo’s. De adoptie van artificiële intelligentie (AI) binnen nieuwsorganisaties verloopt zoals hierboven beschreven (zie 2.4) niet zonder uitdagingen. Terwijl AI-technologieën steeds vaker worden ingezet om journalistieke processen te ondersteunen wijzen experts op een fundamenteel probleem: kennis blijft te vaak opgesloten binnen afzonderlijke afdelingen, waardoor zogeheten silo’s ontstaan die de ontwikkeling en toepassing van AI binnen redacties belemmeren.

Volgens Dodds et al. (2024) vormen kennissilo’s een cruciale belemmering voor de implementatie van ‘responsible AI’ in de journalistiek. Dit benadrukt een bredere uitdaging: AI, ook beeld-AI, wordt soms als een ‘black box’ ervaren, waardoor redacteurs zich onzeker voelen over de technologie en terughoudend zijn om deze kritisch te bevragen of er actief mee aan de slag te gaan.

Naast structurele belemmeringen spelen ook culturele factoren een rol. Een expert in dit onderzoek benadrukte de geslotenheid die binnen redacties kan bestaan als het gaat over beeld-AI-gebruik:

“Ik denk dat er (...) onder journalisten ook een soort geheimzinnigheid over kan zijn omdat er niet per se een open gesprek over wordt gevoerd (...) en ook zeggen van ‘hey heb je dit al gezien? Of dit en dit heb ik geleerd en dit werkt heel goed of juist niet’ en daar dus open over te durven zijn.”

Deze ‘geheimzinnigheid’ kan ervoor zorgen dat redacteuren aarzelen om AI-systemen kritisch te evalueren of best practices met elkaar te delen. Dit is problematisch, omdat AI niet alleen een technologische uitdaging vormt, maar ook een ethische en journalistieke kwestie is. Het gesprek over de ervaringen met AI kan daarentegen op ethisch en journalistiek vlak de ontwikkeling van journalisten juist positief beïnvloeden.

4.4.3 Verlies van gerespecteerd creatief ‘ambacht’ van fotograaf

De opkomst van generatieve AI binnen de journalistieke fotografie leidt tot fundamentele zorgen over de toekomst van het vak. De fotografen onder onze respondenten vrezen niet alleen baanverlies door geautomatiseerde beeldgeneratie, maar wijzen ook op een bredere ondermijning van de geloofwaardigheid van visuele journalistiek. Een ervaren fotojournalist benadrukt deze verschuiving:

“Als mensen een foto in de krant zien of op internet, zien ze het echt als waar. En dat is daardoor heel krachtig. Journalistiek over het algemeen, natuurlijk, maar ook fotojournalistiek. En wat je steeds vaker ziet is dat de nepbeelden niet meer van echt te onderscheiden zijn. En dat raakt eigenlijk ook de fotojournalistiek. Omdat de waarde van beeld, of hoe mensen naar beeld kijken, verandert.”

Daarnaast benadrukken fotografen dat de menselijke factor in beeldjournalistiek niet zomaar kan worden vervangen. Het vak vereist niet alleen technische expertise, maar ook journalistiek inzicht in compositie, context en ethische afwegingen. Zoals een expert het verwoordt:

“Je kan ook niet elke foto in de krant vervangen door een met AI getekende illustratie. Dat kan ook niet. Dus fotografen heb je ook gewoon nodig. Anders dan rollen we van het hellend vlak af, zou ik maar zeggen.”

Deze uitspraak weerspiegelt een bredere angst binnen de sector: wanneer AI-beelden te dominant worden, kan het publiek het vertrouwen in fotojournalistiek verliezen. Hoewel AI dus bij kan dragen aan bepaalde toepassingen, mag dit, volgens onze respondenten, niet ten koste gaan van de kernwaarden van fotojournalistiek: authenticiteit, betrouwbaarheid en menselijk vakmanschap.

Generatieve AI is ten slotte niet creatief in de zin zoals de mens dat is, stellen verschillende respondenten. Hoewel AI een nuttig hulpmiddel is, kan het menselijk vakmanschap en intuïtie bij het creëren van beelden volgens hen niet vervangen. Het unieke vermogen van mensen om emotie, nuance en context in visuele verhalen te brengen, wordt als de hoogste waarde beschouwd door onze respondenten. Dit in tegenstelling tot de eenheidsworst van veel door AI gegenereerde beelden. Een beeldredacteur:

“Een fotograaf kan zijn passie en emotie met een gemaakte foto meer tonen dan een AI-foto. In een AI-foto zit geen diepgang, en daar voel je niet snel een emotie bij.”

4.4.4 Baanverlies

Niet alleen menselijke ambacht komt in het geding door de opkomst van AI: de baan-kansen voor fotografen en beeldredacteurs nemen er ook door af, verwacht een aantal respondenten. Generatieve AI wordt nog maar mondjesmaat ingezet en zal volgens de respondenten niet leiden tot een kaalslag, maar AI-toepassingen om routinematige efficiëntie te vergroten kunnen wel zorgen voor baanverlies. Een respondent zegt:

“Je zult wel iets minder fotografen nodig hebben. We werken bij de AD Regioedacties zonder beeldredacteurs [...] Bij AD Regio centraal zitten straks ook geen vier beeldredacteurs meer. Dat worden er ook twee omdat dingen geautomatiseerd kunnen worden, want iedereen kan in bijna alle databases zoeken. Dus ja, dan gaan ook mensen hun baan verliezen. Dat is de realiteit.”

AI brengt geautomatiseerde, versimpelde en versnelde mogelijkheden om in een database beelden te zoeken. Waar het archiveren, zoeken en selecteren van beelden van oudsher tot de kerntaken van beeldredacteurs behoren, kunnen ‘reguliere’ redacteurs dit door de inzet van AI gemakkelijk zelf.

4.5 Richtlijnen

Welke richtlijnen hebben redacties momenteel voor het omgaan met AI-beelden?

4.5.1 Weinig richtlijnen

Ondanks de toenemende toepassing van AI in journalistieke beeldproductie, ontbreekt het binnen veel mediaorganisaties aan duidelijke en toegankelijke richtlijnen voor het gebruik van beeld-AI. De praktijkexperts die voor dit onderzoek werden geïnterviewd, gaven aan dat er op hun redacties wel gesprekken worden gevoerd over de ethische en praktische implicaties van AI, maar dat officiële richtlijnen grotendeels ontbreken of slechts impliciet aanwezig zijn.

Opvallend is dat de meeste respondenten weinig kennis hadden van bestaande richtlijnen binnen hun eigen organisatie. Degenen die wel richtlijnen kenden, spraken vaak over algemene beleidsdocumenten die niet specifiek voor hun medium waren ontwikkeld, maar eerder concernbrede richtlijnen omvatten, zoals bij DPG Media het geval is.

Bij de regionale publieke omroepen werden er in 2023 eerste kaders gesteld voor AI (met het idee dat richtlijnen later aangescherpt konden worden). Kort gezegd was toen het idee: alleen niet-generatieve AI en altijd een menselijke check. Eind 2024 is de tweede versie van de richtlijnen opgesteld, waarbij ook generatieve AI toegestaan werd.

Een expert benadrukte de urgentie van het vaststellen van duidelijke kaders:

“Ik denk dat die richtlijnen er zeker moeten komen, want de ontwikkelingen gaan zo snel. Je bent al snel te laat en de richtlijn is niet in beton gegoten, dus die kan je altijd nog aanpassen.”

Deze uitspraken sluiten aan bij bredere observaties dat het gebrek aan expliciete richtlijnen en transparantie kan leiden tot onzekerheid bij journalisten over wat wel en niet acceptabel is. Zoals een andere expert opmerkte:

“Als je geen richtlijnen hebt, dan hebben mensen geen enkel referentiekader. Er ontstaan problemen met publieksvertrouwen, omdat niemand weet wat de standaarden zijn of hoe ze deze tools correct moeten gebruiken.”

Hoewel sommige grote nieuwsorganisaties stappen zetten om AI-richtlijnen te ontwikkelen, blijft de implementatie ervan gefragmenteerd. Er is een duidelijk onderscheid tussen richtlijnen voor AI-functionaliteiten (zoals beeldbewerking) en richtlijnen voor AI-generatie (zoals volledige beeldcreatie), maar op veel redacties is het grijze gebied daartussen nog weinig gedefinieerd.

4.5.2 Richtlijnen zijn noodzakelijk

Binnen nieuwsredacties en mediaorganisaties is er een groeiend besef dat duidelijke, werkbare richtlijnen voor AI-gegenereerde beelden noodzakelijk zijn. Terwijl AI steeds vaker wordt ingezet voor beeldcreatie, ontbreekt het volgens de meeste respondenten nog aan uniforme en transparante kaders om de toepassing ervan op een journalistiek verantwoorde manier te reguleren. Dit leidt volgens hen tot onzekerheid binnen redacties en kan het vertrouwen van het publiek in visuele journalistiek ondermijnen. Een expert benadrukte de urgentie van heldere afspraken zo:

“I think it’s very important that media outlets have a charter or at least a flow chart to decide in which cases automatically generated pictures are allowed. And again, this is a discussion that each newsroom should have.”

Samenwerking binnen teams en tussen nieuwsorganisaties wordt als cruciaal gezien voor de ontwikkeling van robuuste richtlijnen. Redacties worden aangemoedigd om niet alleen interne afspraken vast te leggen, maar ook gezamenlijk standaarden te ontwikkelen die aansluiten bij ethische journalistieke principes. Een expert verwoordde deze noodzaak als volgt:

“We zien dat nieuwsredacties er wel steeds meer mee bezig zijn, omdat die evolutie van AI supersnel loopt. [...] Ik denk vooral dat het nu echt dringend tijd is om er echt werk van te maken en dat het er komt.”

Het merendeel van de beeldredacteurs, persfotografen en grafisch vormgevers voelt ook de behoefte om middels richtlijnen met beeld-AI om te gaan. Opvallend is dat veel respondenten nog niet erg op de hoogte zijn van een variëteit aan bestaande richtlijnen, maar daar wel graag meer over willen weten. Een beeldredacteur zegt: “Dat kan voor ons ook heel interessant zijn, want we zijn natuurlijk allemaal zoekende naar een richtlijn die een beetje houdbaar is.” Een andere beeldredacteur bespreekt onderdelen die hij graag aangestipt zou zien in toekomstige richtlijnen:

“Ik denk dat het goed is als dit besproken wordt, van hoe gaan we hiermee om, hoe verhouden we ons tot dit medium binnen verschillende genres in de journalistiek? En ook je beeldrechten: als je eigendomsrecht hebt, wanneer ben jij nog de auteur van het beeld en wanneer niet? Het is wel fijn als daar bepaalde richtlijnen in zitten.”

4.5.3 Richtlijnen op meerdere niveau's

Naast interne richtlijnen is er een dringende behoefte aan sectorbrede en internationale afspraken over het gebruik van AI-gegenereerde beelden in de journalistiek. De snelle technologische ontwikkelingen en de toenemende toepassing van AI binnen nieuwsorganisaties maken het noodzakelijk om journalistieke normen en ethische kaders niet alleen binnen individuele redacties te verankeren, maar ook op nationaal en internationaal niveau af te stemmen. Sommige experts pleiten daarom voor regulering op zowel nationaal als Europees niveau, waarbij richtlijnen worden ontwikkeld die consistent toepasbaar zijn binnen de journalistieke sector. En die zouden helpen om het vertrouwen in beeldjournalistiek te behouden.

Tegelijkertijd benadrukken experts dat de journalistieke sector niet moet wachten op overheidsregulering, maar zelf verantwoordelijkheid moet nemen in het ontwikkelen en implementeren van heldere richtlijnen. Zoals een expert benadrukt:

“I’m not saying it’s too late, but I’m saying that if I was, you know, the major newspapers in the Netherlands or magazines, I would sit down tomorrow and I would say, we have to tell our readers what is what. And we should tell them next week and reassure them and put it on the front page.”

Om tot een effectieve regulering te komen, pleiten deskundigen voor een gelaagde aanpak waarin nationale en wereldwijde standaarden worden gecombineerd met interne redactieprotocollen. Dit betekent dat nieuwsorganisaties niet alleen eigen beleidslijnen moeten ontwikkelen, maar dat ook journalistieke beroepsverenigingen en vakbonden een belangrijke rol kunnen spelen in het opstellen van sectorbrede richtlijnen. De NVJ heeft niet de ambitie gedetailleerde AI-beeldrichtlijnen op te stellen. Zij hebben juist een zeer beknopte lijst van tien richtlijnen voor de gehele journalistiek opgesteld. Ook voor handhaving zien ze geen rol voor zichzelf weggelegd:

“Wie gaat hem of haar daaraan houden? Dat doen wij als beroepsvereniging natuurlijk niet. Wij faciliteren alleen maar, wij oordelen niet.”

Een expert benadrukt de noodzaak van een bredere samenwerking tussen verschillende actoren binnen de media:

“I think supportive peripheral organizations like unions, like journalistic bodies, professional associations can also help in this space. They can develop some guidelines or help organizations develop their own.”

Dit benadrukt de meerwaarde van collectieve afspraken binnen de journalistieke sector, zodat zowel grote als kleinere nieuwsorganisaties toegang hebben tot een gedeeld ethisch kader en uniforme richtlijnen.

4.5.4 Zorgen over externe regelgeving die journalistieke vrijheid en creativiteit beperkt

Alhoewel volgens de geïnterviewde experts de noodzaak bestaat voor duidelijke richtlijnen, roept die externe regelgeving rond AI en beeldgebruik in de journalistiek grote zorgen op bij bepaalde respondenten. Zij vrezen dat dergelijke wetten de persvrijheid en journalistieke creativiteit ernstig kunnen inperken. De fundamentele vraag is: wie bepaalt wat toelaatbaar is? De journalistiek zelf of een externe wetgevende macht? Zoals een van de respondenten het scherp formuleerde:

“Wat wil je kaderen? Wat ga je doen? Ga je de media vertellen waar ze aan moeten voldoen? Ik vind dat de wet zich niet moet bemoeien met wat journalisten doen. Onze taak is de wetgevende macht controleren. Niet andersom.”

Dit sentiment weerspiegelt de bredere zorg dat regulering al snel kan overgaan in censuur, waardoor journalisten minder vrijheid krijgen om AI-technologie op verantwoorde wijze te benutten in hun werk. Een andere respondent waarschuwt expliciet voor deze glijdende schaal:

“Dat wordt al gauw censuur. Daar ben ik wel allergisch voor. Ik zou ook niet goed weten wat voor regels je zou moeten bedenken. Dat lijkt mij ook heel moeilijk. Wat mag wel, wat mag niet.”

De noodzaak om richtlijnen op te stellen en deze actief te implementeren wordt door velen onderstreept. Een respondent stelt:

“Ik denk wel dat elke redactie voor zichzelf daarover nagedacht moet hebben en misschien ook iets vastgelegd moet hebben. In journalistiek is ethiek natuurlijk sowieso een belangrijk punt, van wat laat je wel zien, wat laat je niet zien, wat schrijf je wel?”

4.5.5 Bestaande richtlijnen

Sommige Nederlandse media hebben in hun stijlboeken richtlijnen opgenomen over het gebruik van AI en beeld-AI in hun publicaties. Enkele respondenten verwijzen hiernaar. De richtlijnen gaan onder meer over menselijk toezicht (het human in the loop-principle), ethische standaarden ter voorkoming van fouten en bias, en transparantie over AI-gebruik.

De Volkskrant (oktober 2023) levert eerst een categorische afwijzing:

“De redactie publiceert geen journalistiek werk dat is gegenereerd door kunstmatige intelligentie.”

Later volgt een ander statement:

“Niets wat door AI is gegenereerd, mag onder de naam van de Volkskrant worden gebruikt zonder dat een mens dat goedkeurt.”

ANP (23 april 2023) benadrukt het belang van de eindverantwoordelijke van de mens:

“We publiceren tekst, beeld, graphics of audio die op welke wijze ook door de computer (of door AI) is gegenereerd alleen met toestemming van chefs en hoofdredactie. We zijn daarin zowel onderzoekend als terughoudend omdat AI-teksten er evenzogoed veelbelovend uitzien als fouten kunnen bevatten (‘hallucinaties’) en vooroordelen (‘bias’) kunnen herbevestigen. We gebruiken daarom geen door de computer (of AI) gegenereerde content, ook niet als bronmateriaal, zonder het checken van deze informatie door een mens.”

NU.nl (3 februari 2025) beschrijft in meer detail welke fotobewerkingen acceptabel zijn en dat daarover dan ook transparantie moet zijn:

“NU.nl zal de inhoud van nieuwsfoto’s niet op een onzichtbare manier met AI bewerken. Ook doen we geen inhoudelijke aanpassingen aan nieuwsfoto’s. Alleen kleine, functionele bewerkingen - zoals het contrast aanpassen, croppen of kleurcorrectie - zijn toegestaan.

Als NU.nl meerdere nieuwsfoto’s met behulp van AI samenvoegt, dan moet de bewerking duidelijk zichtbaar zijn. Dat vermelden we ook in het copyright.

Daarnaast is NU.nl terughoudend met door AI gegenereerd beeld. Zo hebben we ethische bezwaren tegen het genereren van gezichten van mensen die niet echt bestaan of situaties die zich nooit hebben voorgedaan.

Maar er zijn ook situaties te bedenken waarin door AI gegenereerd beeld geen bezwaren oproept en meerwaarde heeft. [...] Ook kunnen we met AI beeld genereren voor infographics en ter illustratie van nieuwe AI-toepassingen, maar niet voor stockfoto’s.”

Roularta Media Group [z.d.]:

“Wij zijn niet van plan om zonder redactioneel en menselijk toezicht uitsluitend door AI gegenereerde beelden of video's te publiceren. Beelden gegenereerd door AI worden gebruikt om sommige van onze artikelen te illustreren. Het gaat hierbij om aan de tekst ondersteunende illustraties en geen fictieve foto's. De herkomst van elk AI gegenereerd beeld zal expliciet vermeld worden. Net als bij teksten kunnen we AI gebruiken als hulpmiddel bij het samenvatten van video's, met name om de verspreiding ervan op verschillende platforms te vergemakkelijken.”

Er is dus ruimte voor behoedzaam experimenteren, maar de inzet van beeld-AI bij fotografie in de nieuwsjournalistiek is onbegaanbaar. Inzet bij ondersteunende illustraties is wel mogelijk, maar wel geldt dan het zogenoemde Human in the loop-principe. ANP formuleert het zo in haar AI-leidraad van 2023:

“We houden in onze productieketen vast aan de nu al geldende lijn mens> machine>mens waarbij het bedenken en beslissen bij de mens begint en eindigt, mogelijk in de tussenliggende productie geassisteerd door de computer.”

4.5.6 Concrete richtlijnen van respondenten

Respondenten dragen zelf concrete ideeën voor richtlijnen aan. Voorbeelden van gewenste regels gaan over gedegen beeldverificatie, toegestane en niet-toegestane AI-programma's, en basisregels voor AI en fotobewerking. In dat laatste geval spreken respondenten vaak over bestaande fotobewerkingsregels en -normen, die volgens hen gevormd en uitgebreid kunnen worden. Een grafisch vormgever zegt, bijvoorbeeld: “Ik zou kijken welke regels er nu al zijn voor fotografie en beeldbewerking, en hoe kunnen we die misschien een beetje aanpassen, zodat ze, wat meer bijtijds zijn?” Een fotograaf belicht bovenstaande uitgebreider:

“Waar in de journalistiek wel consensus over is, is in hoeverre je een beeld mag bewerken. [...] Contrast en witbalans is prima, maar echt iets toevoegen of weghalen, ja, dat is niet oké. Sommige persbureaus zijn iets strenger in hoeveel je mag bewerken. Dus de één zegt: we willen het meer hebben zoals het uit de camera komt en de ander zegt ja, je mag wel iets eraan doen. Er is een klein verschil, maar de basis is: je haalt niks uit je foto en je voegt ook niks toe. Zoiets moet je denk ik ook voor AI krijgen: dat je in ieder geval weet wat de basisregels voor AI en fotobewerking zijn.”

4.5.7 De handhaving van richtlijnen

De aantoonbare behoefte aan richtlijnen roept de vragen op wie verantwoordelijkheid moet dragen voor de handhaving ervan, en wat de gevolgen van overtreding van de richtlijnen zijn.

De meeste respondenten vinden niet dat de overheid journalistieke richtlijnen over visuele AI moet handhaven, maar dat collega's of ombudspersonen dat moeten doen. Een ombudsman zegt zelf:

“Mediaorganisaties zijn ontzettend eigenwijs. Ze maken eigenlijk helemaal geen afspraken. Iedereen heeft zijn eigen afspraken. Bij ons, als we er uiteindelijk afspraken over hebben, dan zou je kunnen zeggen dat mensen zelf en de chef van beeldredactie dat moeten handhaven. En als het niet goed gaat, heb je altijd de ombudsman.”

Fotografen tonen doorgaans weinig mededogen als het aankomt op collega's die richtlijnen – of ongeschreven regels – overtreden. Een fotograaf spreekt over een voorbeeld van harde handhaving:

“En er zijn voorbeelden van fotografen, die bijvoorbeeld een rookwolk hebben gekloond. In Beirut, of in ergens in het Midden-Oosten. Nou, die fotograaf was van Reuters. Die was toen dat ontdekt werd, zijn baan kwijt. En daar ben ik het mee eens. Want anders weet de consument niet meer wat de werkelijkheid is. Dat is ook denk ik – nou, dat weet ik zeker - de basis van dat verbod. En ook de basis van de morele drempel waar je niet overheen gaat. Omdat je anders medeplichtig bent aan nepnieuws. En nepnieuws kun je ook in de fotografie creëren.”

Een andere fotograaf benadrukt dat ook vóór het AI-tijdperk er streng werd opgetreden tegen fotografen die beeld manipuleerden, bijvoorbeeld door mensen of dieren in foto's te plakken. “En die mensen worden gewoon ontslagen. Ik vind dat dat een terechte maatregel is.”

4.6 Noodzakelijke AI-expertise

Welke expertise hebben nieuwsredacties nodig in het selecteren, verifiëren en produceren van AI-beeld?

4.6.1 Educatie

Om uitdagingen met betrekking tot transparantie, betrouwbaarheid en ethisch gebruik van AI en nieuwsbeelden effectief aan te pakken, is het essentieel dat zowel journalisten als het bredere publiek worden opgeleid in het herkennen en verantwoord inzetten van AI-content. Dit benadrukt het groeiende belang van onderwijs en mediawijsheidsprogramma's als fundamentele pijlers in de omgang met AI in de journalistiek. Organiseer workshops en trainingen, zeggen de geïnterviewden.

Een expert:

“Bijvoorbeeld workshops te organiseren om journalisten beter vertrouwd te maken met die tools, met die dashboards die we meer toegankelijk willen maken. Zo zouden we bijvoorbeeld op de opleiding journalistiek ook een factcheckingcursus kunnen geven in de toekomst. [...] Of gewoon rechtstreeks op redacties.”

Deze uitspraak illustreert de noodzaak van structurele educatieve initiatieven binnen zowel journalistieke opleidingen als professionele redacties.

Het aanbieden van AI-workshops in factcheckingcursussen en op redacties kan beeld-redacteuren beter voorbereiden op het herkennen en verifiëren van AI-gegenereerde content.

Dit sluit aan bij bredere bestaande initiatieven op het gebied van mediawijsheid (zie bv. Mediawijsheid.nl, Netwerk Mediawijsheid, Mediawijs.be), die niet alleen journalisten, maar ook het bredere publiek in staat stellen kritisch om te gaan met AI-content en de impact ervan op het journalistieke proces te begrijpen.

Uit het beperkte verificatie-experiment van dit onderzoek, waarin respondenten AI-beelden van reguliere foto's moesten onderscheiden, blijkt dat dit onderscheid moeilijk te maken is en dat gestandaardiseerde verificatieprocedures nog zeldzaam zijn.

4.6.2 AI-experts

De ethische en praktische uitdagingen die AI bieden aan de nieuwsmedia maakt de aanwezigheid van AI-specialisten binnen redacties niet slechts wenselijk, maar noodzakelijk, aldus veel respondenten. Deze experts kunnen fungeren als toezichthouders op de implementatie en het gebruik van AI in de journalistieke praktijk, waardoor de integriteit van het vak gewaarborgd blijft.

Een AI-expert binnen de redactie kan bijdragen aan het opstellen en handhaven van duidelijke richtlijnen voor AI-gebruik, het trainen van journalisten in AI-geletterdheid en het bewaken van ethische en juridische kaders. Zoals een expert stelt:

“Ik denk dat elke belanghebbende krant of medium een AI-expert in dienst zou moeten hebben (...) Of in ieder geval iemand die waakt over de procedures. Zodat mensen ook accountable zijn binnen het team.”

Deze uitspraak benadrukt het belang van structurele expertise binnen nieuwsorganisaties. Zonder een dergelijke specialist ontbreekt het aan een centraal aanspreekpunt dat verantwoordelijkheid draagt voor de ethische en praktische toepassing van AI-technologieën in de journalistieke productie. Dit is des te relevanter gezien de razendsnelle ontwikkelingen op AI-gebied, waardoor het voor individuele journalisten onmogelijk wordt om deze technologieën zelfstandig bij te houden (zie ook de bijlage met apps op dit gebied). Een fotograaf uit ons corpus accentueert dit probleem:

“Ja, omdat de ontwikkeling zo snel gaat dat het onmogelijk is om dat in je eentje of van iedereen te verwachten dat ze dat bijhouden. Je bent natuurlijk gewoon al druk met je normale baan. Zeker, freelancers hebben het gewoon vaak druk en geen tijd om ook nog allemaal dat soort dingen bij te gaan houden. En dan is het gewoon superhandig als je iemand hebt die knetterveel weet over AI op de redactie.”

Deze en dergelijke observaties benadrukken dat AI-specialisten niet alleen de technologische kenniskloof binnen redacties overbruggen, maar ook bijdragen aan de werkdrukverlaging van journalisten en freelancers. Het creëren van een dedicated AI-expertisepositie binnen redacties voorkomt dat journalisten alleen zelf de verantwoordelijkheid moeten dragen voor het bijhouden en correct toepassen van AI-tools, terwijl tegelijkertijd de kans vergroot wordt dat AI op een verantwoorde en transparante manier wordt ingezet.

4.7 Transparantie

Welke vormen van transparantie zijn gewenst voor de totstandkoming van het journalistieke beeld, in het bijzonder door AI gegenereerd beeld?

4.7.1 Transparantie essentieel

De opkomst van generatieve AI in de journalistiek heeft geleid tot een groeiende discussie over de noodzaak van transparantie en de implementatie van transparantiemechanismen. Terwijl AI een efficiënte tool kan zijn voor het genereren van visuals, roept het ook vragen op over de geloofwaardigheid en authenticiteit van journalistieke beeldvorming. Het principe van transparantie wordt door vele van onze respondenten gezien als een logisch fundament van journalistiek, net zoals het dat is voor de wetenschap. Zoals een expert het verwoordt:

“I think being transparent is one of the logical conclusions of working to provide the truth, just like for scientists, right? You have to explain your methodology. Otherwise, it’s not science. So I think it’s the same. You have to be transparent. Otherwise, it’s not journalism. But it’s not going to have any effect on trust.”

Deze uitspraak benadrukt dat transparantie niet alleen een ethische verplichting is, maar ook een structureel onderdeel zou moeten zijn van het journalistieke proces (ongeacht het effect op het vertrouwen van het publiek in de journalistiek). Toch blijft de vraag hoe transparantie in de praktijk moet worden toegepast complex en controversieel.

Om de transparantie rond AI-gegenereerde beelden te versterken, worden verschillende mechanismen voorgesteld, waaronder **watermerken** om AI-gegenereerde content visueel te onderscheiden van originele fotojournalistieke beelden. Er is echter discussie over de effectiviteit, omdat watermerken eenvoudig kunnen worden verwijderd of vervalst.

Een andere suggestie zijn **metadatamarkeringen**: embedden van AI-informatie in metadata biedt immers een technologische oplossing die moeilijker te manipuleren is.

Dit vereist echter dat media-instellingen en technologiebedrijven samen standaarden ontwikkelen die breed geaccepteerd worden. Verder worden ook duidelijke *annotaties in beeldcaptions* gesuggereerd. Hoewel het plaatsen van AI-gerelateerde labels in captions transparantie verhoogt, blijft het een uitdaging, omdat uit onderzoek blijkt dat veel lezers captions niet actief waarnemen. Dit benadrukt de noodzaak van aanvullende visuele of technologische markerings.

Een van de praktijkrespondenten wijst erop dat transparantie als concept breed gedragen wordt, maar de implementatie ervan problematisch blijft:

“Transparantie blijft een lastig concept. Iedereen is er groot voorstander van, alleen ‘hoe’ blijft een lastige vraag om te beantwoorden. Transparantie in de caption is niet toereikend genoeg, want weinig mensen lezen die. De oplossing lijkt toch te neigen naar watermerken en AI-informatie in de metadata.”

De discussie over transparantie raakt ook de kern van het vertrouwen in nieuwsmedia. Naarmate AI-beelden realistischer worden, groeit de vrees dat het publiek het onderscheid met authentieke fotografie verliest. Dit kan leiden tot een erosie van vertrouwen in journalistiek als instituut. De vraag blijft in hoeverre transparantiemaatregelen daadwerkelijk effect hebben op het versterken van publiek vertrouwen, of dat er juist een afnemend geloof in journalistieke fotografie ontstaat door de groeiende aanwezigheid van AI-gegenereerde beelden.

Een klein aantal respondenten filosofeert over vormen van transparantie die ‘de andere kant op’ gaan: bijvoorbeeld het aanpassen van auteursrecht aan AI, in plaats van te verwachten dat AI ooit zal passen in het huidige auteursrecht. Een expert stipt dit kort aan: “Misschien moeten ons hele beeld van auteursrechten wel op de schop. Is het juist andersom? Hebben we een rechtssysteem dat niet werkt?”

Een andere expert pleit voor ‘omgekeerde’ transparantie in de vorm van duidelijkheid over wanneer een beeld authentiek – geschoten met een camera – is, in plaats van duidelijkheid over wanneer een beeld (met AI) gemanipuleerd is. Hij vergelijkt dit met de universeel geaccepteerde, welbekende aanhalingstekens die authenticiteit van tekst beloven:

“If you quote somebody in a text, you know you can’t add words. You can’t change it. So with a photograph, if you think of the frame of the photograph as quotation marks, then maybe you shouldn’t be able to change in it just the way it quotes for words. Whatever is inside the frame is like quotation marks. It’s a very easy concept. If we do that and we tell readers this is the way we work with photographs, we pretend that the frame is like a quotation mark. We don’t change what’s in it except for tiny things like a little bit of contrast or something like that. The reader understands. These are simple ideas.”

4.7.2 Aarzeling over ‘te’ transparant zijn

Uit onze interviews blijkt dat zowel experts als journalisten aarzelen over de mate van transparantie die nodig is bij het gebruik van generatieve AI in nieuwsbeelden. Hoewel transparantie een fundamenteel principe is binnen de journalistiek, leeft de zorg dat een te grote nadruk hierop juist averechts kan werken. Te veel informatie over de inzet van AI zou het publiek kunnen verwarren, hen afkeren van nieuwsmedia of zelfs het vertrouwen in journalistieke inhoud verminderen. Deze paradox plaatst nieuwsorganisaties voor een complexe uitdaging: hoe kan men enerzijds open zijn over AI-gebruik, zonder dat dit leidt tot onterechte scepsis of informatie-overload?

Hoewel een gebrek aan transparantie journalistiek minder betrouwbaar maakt volgens de interviewers, stellen veel van hen dat transparantie op zichzelf geen garantie biedt voor vertrouwen. Zoals een respondent meldt: “Transparantie op zichzelf heeft geen effect op vertrouwen.” Een van de meest gehoorde zorgen is dat het publiek de extra informatie over AI-gebruik in nieuwsbeelden niet leest of begrijpt, waardoor de intentie van transparantie zijn doel voorbijschiet. Een respondent verwoordde dit als volgt:

“Als je enkel een link toevoegt, lijkt het alsof je iets verbergt. Maar tegelijkertijd hoeft niet alles gelezen te worden. En hoe vaak wordt die extra uitleg eigenlijk gelezen? Klikt iemand er überhaupt op?”

Dit roept vragen op over de effectiviteit van transparantiemechanismen zoals metadata, labels of verklarende teksten bij beelden. Zelfs als AI-gebruik expliciet wordt vermeld, betekent dit niet dat het publiek deze informatie daadwerkelijk leest of erop vertrouwt.

Daarbij komt nog een extra risico: zelfs wanneer AI-transparantie wordt aangeboden, verdwijnt deze informatie vaak uit het collectieve geheugen. Zoals een andere respondent stelt:

“Ik zeg altijd: beeld moet je zonder tekst kunnen verklaren. Er blijft dus aan het beeld hangen dat het waar is, want je hebt het gezien, het staat geschreven, dus het is waar. Je hebt het gezien, dus het is waar. Dus uiteindelijk, als je over vijf jaar kijkt, stel dat het een heel heftig beeld is. Dan wordt er niet meer gezien: ja maar dat was toen AI. Mensen zijn niet zo genuanceerd in hun geheugen. Als het een tekening was geweest of een schilderij, dan zien mensen dat het een afbeelding van de werkelijkheid is.”

Deze uitspraak benadrukt dat het risico van AI-gebruik in journalistieke beelden verder reikt dan het moment van publicatie. In de toekomst kunnen AI-gegenereerde beelden losgekoppeld raken van hun oorspronkelijke context, waardoor het onderscheid tussen echte en kunstmatige content vervaagt (“Het was toch te zien bij NOS”). Dit versterkt de noodzaak voor duurzame transparantiemaatregelen die niet alleen het moment van consumptie, maar ook de langdurige archivering en herkenbaarheid van AI-beelden adresseren.

Naast de discussie over het doel, de vorm en de effectiviteit van transparantie, vragen sommige respondenten zich af wanneer transparantie op zijn plaats is. Is het noodzakelijk om AI-gebruik expliciet te vermelden wanneer AI slechts een minimale rol speelt, bijvoorbeeld bij kleurcorrectie of het verwijderen van stofdeeltjes? Of ontstaat de noodzaak voor transparantie pas wanneer AI een volledige afbeelding of video genereert? Een grafisch vormgever vergelijkt dit met tekstverwerking:

“Met beelden zijn mensen vaak heel kritisch: bij een video vindt iedereen dat erbij moet staan of iets AI-gegenereerd is, terwijl niemand bij elke nieuwskop een sterretje zet dat het door AI is gemaakt.”

Deze opmerking plaatst de vraag naar AI-transparantie in een breder perspectief: in hoeverre is AI-gebruik bij beeldbewerking wezenlijk anders dan AI-gebruik bij tekstproductie? Bovenstaande discussie roept ook vragen op over de status van beeldbewerking met Adobe Photoshop, een programma dat door veel fotografen en beeldbewerkers als standaard wordt gebruikt. Photoshop bevat steeds meer AI-functionaliteiten, zowel in de vorm van generatieve AI via Adobe Firefly als in meer subtiele toepassingen zoals automatische kleurcorrectie, het selecteren van objecten of het croppen van afbeeldingen. Dit leidt tot een fundamentele vraag bij de respondenten:

“Wat gaan we dan labelen? [...] Bij sommige tools zit het gewoon ingebakken: in Photoshop zit AI. [...] In theorie zou je op den duur alles kunnen labelen. Moet je iedere aanpassing die met AI gemaakt werd, of waarin AI een hulpje was, gaan aanduiden? Ik denk niet dat labelen zo ver moet gaan.”

Deze uitspraak illustreert dat het debat over AI-transparantie niet alleen draait om generatieve AI, maar ook over bredere automatisering binnen beeldbewerking. Wanneer AI diep geïntegreerd is in journalistieke productiemiddelen, zoals Photoshop, Premiere Pro of geautomatiseerde ondertitelingssoftware, wordt het steeds moeilijker om een duidelijke grens te trekken tussen AI-geassisteerde en AI-gegenereerde content.

4.8 De waarheid van beelden in de journalistiek

Filosofische reflecties op tekst, beeld en waarheid in de journalistiek benadrukken de verschillende percepties van waarheid die aan deze media worden toegeschreven. Foto's worden vaak als waarheidsgetrouw beschouwd vanwege hun systematische verwijzing naar de werkelijkheid en hun mimetische kwaliteiten, zoals beschreven door Arends & De Jong (2024). De aard van fotografie als medium speelt een cruciale rol in de geloofwaardigheid ervan, vooral in vergelijking met tekst, dat van nature meer interpretatief en symbolisch is. Waar tekst wordt gezien als een 'code' die door interpretatie en bewerking kan worden gevormd, wordt fotografie doorgaans beschouwd als een objectieve weergave van de werkelijkheid. Zoals door een van de experts in onze interviews wordt gesteld:

“Fotografie wordt gezien als een drager van de waarheid met een hoofdletter W, terwijl tekst inherent interpretatief is. Mensen verwachten dat tekst wordt geredigeerd of zelfs gekleurd, maar foto's worden verondersteld de realiteit te tonen zoals deze is.”

“Fotografie wordt gezien als een drager van de waarheid met een hoofdletter W, terwijl tekst inherent interpretatief is. Mensen verwachten dat tekst wordt geredigeerd of zelfs gekleurd, maar foto’s worden verondersteld de realiteit te tonen zoals deze is.”

Dit (perceptie)verschil werpt belangrijke vragen op over hoe beide media (tekst en beeld) bijdragen aan de constructie van waarheid in de journalistiek. Is fotografie minder subjectief dan geschreven journalistiek? Of is het ‘redigeren’ binnen de fotojournalistiek (door encenering, kadering, lenskeuze, beeldselectie, bijsnijden, kleurbewerking) enkel minder transparant en dus mogelijk juist meer misleidend dan in de geschreven journalistiek? Is een verminderd geloof in de objectiviteit van fotografie een kwalijk gevolg van Photoshop en AI, of juist een logische en gezonde correctie op een hardnekkige misvatting de foto als unieke waarheidsdrager? In welke situaties kunnen persfoto’s nog een unieke bewijsfunctie vervullen?

Hoofdstuk 5

Conclusies en aanbevelingen

Conclusies en aanbevelingen

5.1 Conclusies

Op basis van de resultaten van dit onderzoek kunnen we verschillende conclusies trekken over de rol en impact van AI-gegenereerde beelden in de Nederlandse journalistiek. Onderstaande conclusies komen voort uit aangetroffen uitkomsten (5.1.1 tot en met 5.1.4) en uit risico's van beeld-AI voor de journalistiek (5.1.5 tot en met 5.1.9).

5.1.1 AI-beeldtoepassingen worden nog slechts beperkt en voorzichtig ingezet

De respondenten konden slechts een beperkt aantal 'geslaagde' voorbeelden noemen van journalistieke beeldproductie met AI. Veel toepassingen komen uit het buitenland, terwijl Nederlandse redacties nog zoekende zijn naar hoe AI op een verantwoorde manier kan worden gebruikt. De experimenten die wel zijn gedaan, zoals de AI-nieuwslezer Nina bij Omroep Brabant, stuiten op ethische en praktische bezwaren. De beperkte voorbeelden wijzen erop dat er nog geen breed gedeeld referentiekader bestaat voor succesvolle AI-toepassingen in journalistieke beeldproductie.

Een domein waarin media wel vaak gebruik maakt van beeld-AI is illustratie, dus afbeeldingen zonder pertinente werkelijkheidspretentie. Ook in de afdelingen satire en vermaak (columns) wordt beeld-AI ingezet, maar die toepassingen worden nadrukkelijk onderscheiden van de journalistiek.

5.1.2 AI wordt vooral gebruikt voor routinematige efficiëntie en beeldbewerking

AI blijkt vooral nuttig bij het versnellen van beeldzoek- en bewerkingstaken en het verbeteren van workflow-processen binnen de beeldredactie. Zoek-, selectietools en beeldbewerkingsmogelijkheden worden door beeldredacteuren en fotografen gewaardeerd omdat ze repetitieve taken kunnen automatiseren. AI wordt hier vooral ingezet om basale correcties sneller uit te voeren, waardoor professionals meer tijd hebben voor inhoudelijke en creatieve taken.

5.1.3 AI heeft potentie als tool voor illustratieve en abstracte journalistieke visuele toepassingen

AI wordt momenteel vooral ingezet voor het creëren van illustraties bij opiniestukken of conceptuele onderwerpen zoals inflatie of klimaatverandering. Het vermogen om snel aangepaste stockfoto's, infographics en alternatieve visualisaties te genereren, wordt gezien als een meerwaarde voor de journalistieke praktijk. Toch blijft de toepassing van AI voor pure nieuws- en reportagefotografie zeer omstreden.

5.1.4 AI maakt het mogelijk om ‘onmogelijke’ beelden te genereren

Voor bepaalde situaties waarin het niet haalbaar is om traditionele foto's te maken, zoals door privacybeperkingen, gevaren in conflictgebieden of historische reconstructies, biedt AI nieuwe mogelijkheden. AI kan bijvoorbeeld helpen bij het creëren van visuele representaties van gebeurtenissen uit het verleden of van toekomstprojecties die anders lastig te verbeelden zijn.

5.1.5 Angst voor baanverlies onder beeldprofessionals

Hoewel AI vooral als ondersteunend instrument wordt gepresenteerd, leeft onder beeldmakers en fotografen de vrees dat AI de noodzaak voor menselijke fotografen en beeldredacteurs zal verminderen. Net als bij eerdere digitaliseringsgolven in de journalistiek bestaat de zorg dat AI wordt ingezet als kostenbesparend alternatief, ten koste van vakmanschap en authenticiteit.

5.1.6 Journalistieke richtlijnen voor AI-beeldgebruik zijn nog onvoldoende ontwikkeld

Veel redacties hebben nog geen duidelijke richtlijnen over hoe AI-gegenereerde beelden verantwoord kunnen worden gebruikt. Daar waar richtlijnen bestaan, zijn deze vaak versnipperd, niet gestandaardiseerd en vaak onbekend bij de ondervraagde beeldredacteurs en fotojournalisten. De journalistieke beroepsgroep bevindt zich in een fase van experimenteren, meestal zonder vaststelling van heldere ethische of praktische kaders.

5.1.7 Er is een gebrek aan AI-expertise en kennisdeling binnen redacties

Op enkele gunstige uitzonderingen na, bijvoorbeeld de beeldredactie van ANP, blijkt dat veel beeldredacteurs en fotografen beperkte kennis hebben over hoe AI-tools precies werken en over het detecteren van AI-beeld. Dit gebrek aan expertise kan leiden tot ongewenst gebruik of misinterpretatie van AI-beelden.

Uit het verificatie-experiment van dit onderzoek, waarin respondenten AI-beelden van reguliere foto's moesten onderscheiden, blijkt dat dit onderscheid steeds moeilijker te maken is en dat gestandaardiseerde verificatieprocedures nog zeldzaam zijn.

Daarnaast blijkt er sprake te zijn van 'kennissilo's' binnen redacties, waarbij AI-expertise niet gelijkmatig wordt verdeeld over teams en afhankelijk blijft van een zeer kleine groep specialisten.

5.1.8 Transparantie over AI-gebruik is noodzakelijk maar uitdagend

Hoewel transparantie over AI-gebruik wordt gezien als een onontkoombare journalistieke verplichting, vermoeden de respondenten dat expliciete labeling van AI-beelden soms juist wantrouwen bij het publiek vergroot. Dit ‘transparantiedilemma’ vraagt om een zorgvuldige afweging over hoe AI-gebruik wordt gecommuniceerd. Enerzijds is er een legitieme behoefte aan transparantie om de betrouwbaarheid van journalistiek te waarborgen, maar anderzijds bestaat het risico dat te veel of onhandig gepresenteerde informatie het publiek verwart of afschrikt.

De respondenten verwijzen naar experimenten met subtiële vermeldingstechnieken, zoals emblemen of disclaimers, om het publiek op een genuanceerde manier te informeren, zonder onnodige verwarring te veroorzaken. Dit wijst erop dat transparantie meer is dan een technisch of beleidsmatig vraagstuk; het is een communicatief dilemma dat een doordachte balans vereist tussen toegankelijkheid en begrijpelijkheid. Bovendien wordt AI steeds vaker ingezet als een subtiële assistent binnen bewerkingssoftware, wat vragen oproept over de grenzen van transparantie: moet elk gebruik van AI vermeld worden, zelfs als het slechts kleine optimalisaties betreft? De uitdaging voor redacties is dan ook om een vorm van transparantie te ontwikkelen die het vertrouwen van het publiek versterkt, zonder AI-gebruik onnodig te problematiseren.

5.1.9 AI roept fundamentele vragen op over waarheid en de rol van visuele journalistiek

De toename van AI-beelden in de journalistiek zorgt voor een diepgaand debat over objectiviteit en waarheid. De grens tussen ‘echte’ en gegenereerde beelden vervaagt, wat het risico op misleiding en wantrouwen in de journalistiek vergroot. Redacties moeten zich daarom niet alleen bezighouden met praktische implementatie, maar ook met bredere ethische vragen over de rol en impact van AI op visuele journalistiek. Een gelijkwaardig debat werd gevoerd na de introductie en democratisering van digitale beeldbewerkingsprogramma’s als Photoshop. Dat zette verschillende binnen- en buitenlandse mediaorganisaties ertoe aan om hun gebruik van deze technologie vast te leggen in publieke richtlijndocumenten.

Maar de alomtegenwoordigheid van steeds overtuigendere AI-gegenereerde nepfoto’s brengt ook de bewijswaarde van de foto in het algemeen in gevaar. Die beelden zijn zelfs met geavanceerde detectietools niet altijd van echte foto’s te onderscheiden. Redacties moeten dus, meer dan al het geval was, terughoudend omgaan met het gebruik van foto’s uit ‘onbekende’ bron, zoals burgers of amateurs, als bewijsmateriaal.

5.1.10 AI reproduceert vooroordelen

AI-modellen trainen zich op materiaal waarin onvermijdelijk vooroordelen voorkomen, waardoor de eindproducten van deze modellen ook stereotyperend kunnen zijn. De journalistieke waarde om nieuws te verspreiden waarin mensen ongeacht huidskleur, seksualiteit en geslacht gelijkwaardig aan elkaar zijn, kan door stereotyperende AI worden bedreigd.

5.2 Aanbevelingen

Op basis van dit onderzoek worden de volgende aanbevelingen geformuleerd.

5.2.1 Ontwikkel en implementeer AI-richtlijnen

Om AI op een verantwoorde manier in de journalistieke praktijk te integreren, is het essentieel dat redacties **duidelijke redactionele kaders** vaststellen. AI moet primair worden beschouwd als een hulpmiddel ter ondersteuning van beeldproductie, niet als een vervanging van journalistieke fotografie of illustratie. Redacties leggen expliciet vast in welke situaties AI-gebruik acceptabel is en wanneer het niet toegepast wordt. Heldere richtlijnen helpen niet alleen om ethische grenzen te bewaken, maar zorgen er ook voor dat AI-inzet en communicatie over het gebruik transparant en consistent plaatsvindt.

Een belangrijk **onderscheid** dat hierbij gemaakt kan worden, is dat **tussen illustratieve en feitelijke toepassingen**. AI-gegenereerde beelden kunnen een waardevolle rol spelen in opiniestukken, infographics en abstracte visualisaties, waar het beeld primair een ondersteunende of conceptuele functie heeft. In de context van journalistieke verslaggeving, waarin beeldmateriaal een documentair karakter heeft en een getrouwe weergave van de werkelijkheid ('feiten') moet bieden, is het gebruik van AI-gegenereerde beelden echter problematisch en onwenselijk. Dit onderscheid wordt expliciet in redactionele richtlijnen vastgelegd om misleiding te voorkomen en het vertrouwen in journalistiek beeldmateriaal te waarborgen.

Daarnaast is het cruciaal dat AI-beeldcreatie altijd onder menselijk toezicht plaatsvindt. Het **'Human in the Loop'-principe** (mens-machine-mens) is een standaardpraktijk binnen nieuwsorganisaties: AI mag beelden genereren of bewerken, maar de uiteindelijke beslissing over publicatie ligt altijd bij een menselijke redacteur. Dit voorkomt dat ongecontroleerde AI-algoritmen de beeldselectie en -productie domineren en garandeert dat journalistieke normen en ethische overwegingen worden meegenomen in de eindbeslissing.

Maar dit principe is geen panacee. Hoge werkdruk, een ogenschijnlijk betrouwbaar AI-model, en moeilijk detecteerbare fouten, brengen het risico met zich mee dat de rol van de menselijke redacteur verwordt tot wat AI-experts van Technische Universiteit Delft een 'morele kreukelzone' noemen: een situatie waarbij vrijwel al het werk geautomatiseerd is, maar de human-in-the-loop de morele verantwoordelijkheid op zich neemt wanneer er iets misgaat.

5.2.2 Bied transparantie op maat

Een effectief transparantiebeleid vereist een gedifferentieerde aanpak waarin de mate en het doel van AI-gebruik bepalen hoe expliciet dit wordt gecommuniceerd. Kleine bewerkingen, zoals kleuraanpassingen of ruisverwijdering, hoeven niet expliciet vermeld te worden, terwijl volledig AI-gegenereerde journalistieke beelden een duidelijke duiding vereisen. Daarnaast denken redacties na over de langetermijnherkenbaarheid van AI-beelden: hoe kan worden gewaarborgd dat content ook na jaren nog als AI-geproduceerd wordt herkend? Dit vraagt om structurele markeringen, zoals metadata of subtiele visuele labels.

Het journalistieke veld waarborgt transparantie zonder onnodig wantrouwen te creëren. In journalistieke beeldproductie betekent dit dat redacties duidelijke en reproduceerbare methodes hanteren voor het documenteren van hoe AI wordt ingezet. Dit vraagt een sectorbrede discussie over wat, hoe en wanneer AI-gebruik expliciet benoemd wordt. Transparantie duidt niet alleen de **herkomst van beelden**, maar ook de **mate van AI-bewerking**: eenvoudige beeldoptimalisaties vragen minder strikte maatregelen dan AI-beelden die journalistieke realiteit nabootsen. Hoe AI-gebruik wordt gecommuniceerd, is minstens zo belangrijk: een combinatie van visuele cues, metadata en redactionele verantwoording kan een evenwichtige oplossing bieden. Wanneer transparantie nodig is, hangt af van de impact op beeldintegriteit. In gevoelige journalistieke contexten, zoals oorlogsverslaggeving en politieke berichtgeving, is maximale transparantie vereist, terwijl voor illustratieve toepassingen subtielere duiding volstaat.

Ook al is wetenschappelijk bewijs dat vertrouwen in de journalistiek versterkt wordt door transparantiemaatregelen niet eenduidig, het bieden van transparantie blijft een journalistieke plicht. Door AI-transparantie contextafhankelijk te benaderen en structurele oplossingen implementeren, kunnen redacties het vertrouwen van hun publiek behouden zonder AI-gebruik onnodig te problematiseren.

5.2.3 Wees kritisch op door AI gereproduceerde vooroordelen

Bij gebruik van generatieve AI-modellen wordt rekening gehouden met bias, oftewel de reproductie van vooroordelen. Om de verspreiding van stereotypering en discriminerende afbeeldingen te voorkomen, doen journalisten er goed aan om van tevoren te bedenken welke afbeeldingen (niet) wenselijk zijn bij de verspreiding van nieuws. Bij gebruik van een AI-beeldgenerator wordt de specifieke beeldwens aandachtig en gedetailleerd geprompt om te voorkomen dat AI ontbrekende instructies zelf opvult met stereotypisch of onjuist materiaal.

5.2.4 Versterk de controle op AI-beeldmanipulatie

Eerder onderzoek en de resultaten van dit onderzoek wijzen uit dat het vaststellen van beeldmanipulatie of beeldgeneratie met het blote oog niet meer volstaat. Redacties hebben geen vaststaande procedures om te controleren of aangeleverde beelden authentiek zijn of niet.

Zij vertrouwen erop dat de hun bekende fotografen en persbureaus echte, niet-gemanipuleerde foto's aanleveren. Echter, bij beeld dat via onbekende bronnen en social media aangeboden wordt – en zeker Fin crisissituaties –, zijn gestandaardiseerde procedures die beeldmanipulatie controleren geen overbodige luxe. Er is controle nodig om weerstand te bieden tegen het groeiende aantal gegenereerde en gemanipuleerde afbeeldingen dat circuleert en ook bij nieuwsredacties terechtkomt.

5.2.5 Versterk AI-expertise op de redactie

Om redactiemedewerkers in staat te stellen AI-beeldgeneratie en -detectie kritisch te beoordelen, is **gerichte kennisontwikkeling** essentieel. Nieuwsorganisaties investeren in training en educatie, zoals gebeurt bij DPG Campus, het interne opleidingsplatform van mediaconcern DPG, en bij de NPO, waar interactieve sessies worden georganiseerd om journalisten te informeren over AI-tools en hun journalistieke meerwaarde. Bij RPO is een AI roadshow langs alle omroepen gedaan met een lancering van richtlijnen en een AI toolkit, en er worden regelmatig AI-update-calls georganiseerd om kennis uit te wisselen.

Multidisciplinaire samenwerking tussen journalisten en AI-experts is cruciaal. Door gespecialiseerde teams op te zetten waarin technische en redactionele kennis samenkomen, kunnen redacties AI op een verantwoorde manier integreren in hun werkproces. Bijvoorbeeld bij de RPO is binnen het transformatieteam zo'n multidisciplinair AI-projectteam opgezet. Van hieruit worden de omroepbrede AI-initiatieven genomen, tools aangekocht of getest, kennis gedeeld en richtlijnen voorgesteld aan de hoofdredacteurs.

Redacties zouden moeten beschikken over **betrouwbare beelddetectietools** om AI-gegenereerde content effectief te herkennen en te verifiëren. Enkele nieuwsredacties onderzoeken bijvoorbeeld technologieën zoals de Coalition for Content Provenance and Authenticity (C2PA) om de herkomst en authenticiteit van beelden te verifiëren (Thomson et al, 2025). Investeren in zowel opleidingen als technologische ondersteuning zorgt ervoor dat redacties AI niet alleen benutten als hulpmiddel, maar ook grip houden op de risico's ervan.

Wat is de klimaatimpact van AI? Redacties creëren bewustzijn **over potentiële impact van de klimaatimpact** van AI-beeldgeneratie. AI-tools verbruiken veel energie en redacties kiezen bewust voor energie-efficiënte modellen en duurzame datacentra.

5.2.6 Train en informeer redacties over AI-gebruik

Trainingen over AI-gebruik kunnen lacunes in huidige kennis van beeldredacteurs en fotojournalisten opvullen. Lacunes zijn onder andere aanwezig in **technische expertise** over de mogelijkheden van beeld-AI. Training in het schrijven van doeltreffende prompts gewenst. Die training zou naast technische vaardigheden, ook basiskennis over de aard en gevolgen van maatschappelijke stereotypes behandelen.

Een groot deel van de respondenten erkent dat beeld-AI **ethische bezwaren** met zich meedraagt en maakt keuzes daarover op basis van intuïtie of abstracte richtlijnen. Een ethisch knelpunt is bijvoorbeeld het gebruik van generatieve AI-modellen die getraind worden op basis van onethisch verkregen data. Opleidingen bij redacties kunnen inzicht bieden in hoe specifieke AI-modellen ‘scoren’ op ethisch gebied. Bijvoorbeeld: zijn de modellen transparant over het gebruikte trainingsmateriaal (en is dat materiaal gebruikt met toestemming van de makers?), hoe zit het met de manier waarop het model omgaat met het reproduceren van vooroordelen, en onder welke arbeidsomstandigheden werken de mensen die het model continu optimaliseren? Zulke kennis van is voor beeldredacteurs en fotojournalisten nuttig om weloverwogen te kiezen voor een applicatie.

Om journalistieke betrouwbaarheid te kunnen blijven waarborgen worden beeldredacteurs opgeleid in het vaststellen of een beeld gemanipuleerd of gegenereerd is. Daarbij kunnen specifieke **digitale detectietools** worden ingezet. De gebruiksvriendelijkheid van deze tools verschilt sterk, en sommige AI-detectietools zijn alleen bruikbaar voor zeer gespecialiseerde AI-experts. Bovendien is geen enkel AI-detectieprogramma 100% waterdicht. Het is dan ook te verwachten dat naarmate beeldgeneratiemodellen geavanceerder worden, steeds meer AI-beelden ‘ondetecteerbaar’ worden.

Voorts moeten dit soort trainingen ook in het **journalistieke onderwijs** op hogescholen en universiteiten ontwikkeld en verzorgd worden, zodat AI een structureel onderdeel van het curriculum wordt.

5.2.7 Experimenteer, maar met duidelijke grenzen

De snelle opkomst van AI-beeldgeneratie vraagt om een experimenteercultuur binnen redacties die zowel innovatief als kritisch is. AI-experimenten helpen immers om de mogelijkheden en beperkingen van de technologie beter te begrijpen, mits ze duidelijk begrensd en goed gedocumenteerd worden om de journalistieke geloofwaardigheid te waarborgen.

Redacties testen AI het best in een **gecontroleerde omgeving**, zoals pilotprojecten voor illustraties bij achtergrondverhalen en opiniestukken of historische visualisaties, het animeren van zowel historische als hedendaagse beelden of om data te visualiseren, bijvoorbeeld door van spreadsheets afbeeldingen te maken. **Evaluatie** is hierbij echter cruciaal: interne toetsing bepaalt welke toepassingen journalistiek verantwoord zijn en welke risico's ze met zich meebrengen. Een ‘sandbox’, een afgeschermd digitale testomgeving, kan helpen om tools veilig uit te testen zonder directe impact op het nieuwsaanbod.

Daarnaast is structurele evaluatie en **kennisdeling** van belang. Redacties bespreken AI-experimenten in redactievergaderingen en delen best practices met vakorganisaties. Indien AI-gegenereerde beelden worden ingezet, wordt het publiek hierover transparant geïnformeerd, bijvoorbeeld via een vast format waarin wordt toegelicht wanneer en waarom AI is gebruikt. Zo blijft AI een gecontroleerd hulpmiddel binnen een verantwoorde journalistieke praktijk.

5.2.8 Betrek het publiek in de AI-discussie

Alhoewel AI een groeiende rol speelt in journalistieke beeldproductie, zijn veel mediagebruikers zich niet bewust van de impact ervan. Om vertrouwen in AI-gebruik te versterken, kunnen nieuwsorganisaties transparant communiceren en mediawijsheid bevorderen. Dit kan door, bijvoorbeeld, toegankelijke uitleg over hoe AI wordt ingezet, bijvoorbeeld via vaste rubrieken of korte explainers, zodat het publiek begrijpt welke keuzes redacties maken. Ook het organiseren van AI-webinars en educatieve campagnes kan het lezerspubliek beter informeren over bepaalde journalistieke beslissingen rond AI en beeld.

Daarnaast kunnen media het kritisch nadenken over AI-beelden stimuleren. Samenwerking met onder meer onderwijsinstellingen en factcheckorganisaties kan helpen om het publiek zich bewust te maken van hoe synthetische beelden werken en hoe ze te herkennen zijn. Open dialoog is daarbij essentieel: redacties kunnen lezers actief betrekken bij ethische afwegingen door middel van publiekspeilingen, debatten of feedbackkanalen. Ook columns van bijvoorbeeld ombudspersonen kunnen hieraan bijdragen.

5.2.9 Aanzet tot concrete richtlijnen

Op basis van deze aanbevelingen, zijn de volgende richtlijnen aanbevelen voor een verantwoorde en ethische omgang met AI in de journalistieke beeldproductie.

1. Behoud de menselijke eindverantwoordelijkheid

Journalisten behouden altijd de controle over AI-gegenereerd beeld en de eindbeslissing nemen over de publicatie ervan.

2. Wees transparantie over AI-gebruik

AI-gegenereerde beelden worden expliciet gelabeld, zodat het publiek begrijpt hoe ze tot stand zijn gekomen. Dit kan door watermerken, metadata of een korte uitleg in het artikel.

3. Vermijd misleiding en versterking van stereotypen

AI-beeld wekt niet de indruk dat het een foto is van een echte gebeurtenis en bevestigt geen schadelijke vooroordelen. Journalisten prompten bewust om mogelijke bias te vermijden.

4. Gebruik geen AI-gegenereerde beelden bij feitelijke nieuwsverslaggeving

AI mag worden gebruikt voor illustraties, infographics en visualisaties, maar niet voor beelden die een echt nieuwsfeit weergeven.

5. Respecteer auteursrechten

Gebruik enkel AI-tools die getraind zijn op gelicentieerd of open-source materiaal. Het is cruciaal om te weten of de brongegevens legaal zijn verkregen.

Hoe kan de journalistiek op een verantwoorde manier omgaan met de toenemende mogelijkheden en risico's van beeldbewerking en -productie met AI?

Mogelijkheden

Nieuwe manieren van beeldproductie en tijdsbesparing bij onder meer:

- Infographics
- Visualisaties van niet-fotografeerbare zaken
- Satirische beelden
- Illustraties
- Illustratieve foto's

Risico's

- **Reproductie van stereotypen** in gegenereerde afbeeldingen
- **Kwaliteitsdating:** vervanging van genuanceerd mensgemaakt beeld door AI-eenhedsworst.
- **Schending van beeldrechten** door 'gestolen' trainingsdata.
- **Verlies van de bewijswaarde** van foto's
- **Baanverlies** bij (foto)journalisten

Daarnaast heeft AI-beeldgeneratie een klimaatimpact, en gebeurt de training van de meest bekende AI-modellen door 'data labellers' die werken in slechte arbeidsomstandigheden.

De geïnterviewde experts pleiten naast AI-beeldrichtlijnen op redactieniveau voor **sectorbrede regulering** door beroepsverenigingen, vakbonden en raden voor journalistiek. Beeldmakers, ombudspersonen en redacteurs benadrukken tegelijk het belang van de **journalistieke vrijheid**.

Vijf basisprincipes voor verantwoord journalistiek AI-gebruik*

1. Journalisten **houden het heft in handen** en hebben eindverantwoordelijkheid over beeldproductie (mens-machine-mensprincipe)
2. Journalisten **zijn transparant over AI-gebruik**, maar waken voor een 'informatie-overload'
Hoe uitgebreid en nadrukkelijk die transparantie moet worden doorgevoerd, hangt af van het journalistieke genre en de specifieke beeldtoepassing (zie #4)
3. Journalisten **behoeden zich voor de reproductie van stereotypen** in AI-beeld
Hiervoor is het belangrijk dat redacties investeren in de opleiding van journalisten, zowel op technisch gebied (hoe laat je AI doen wat je wilt) als over beeldvorming en stereotypen.
4. Journalisten maken een duidelijk **onderscheid tussen illustratief beeld en 'feitelijk beeld'**
Om het vertrouwen in de journalistieke foto als 'bewijs' voor echte situaties te beschermen, is het volgens beeldredacteurs belangrijk dat er genrespecifieke normen worden gehanteerd:

Journalistiek genre	Beeldgeneratie	AI-fotobewerking
Satirische beelden	Geen extra beperkingen	Geen extra beperkingen
Niet-fotorealistische illustraties	Geen extra beperkingen	Niet van toepassing
Fotorealistische illustraties	Grote vrijheid, mits extra transparantiemaatregelen om misleiding te voorkomen	Niet van toepassing
Illustratieve foto's	Nooit toegestaan. AI-beelden mogen niet als 'foto's' gepresenteerd worden	Naast belichting en kleurcorrectie mogen in sommige gevallen ook elementen worden weggehaald
Nieuwsfoto's	Nooit toegestaan.	Alleen voor belichting, kleurcorrectie en het weghalen van lensstofjes

5. Journalisten **respecteren de rechten van beeldmakers**
Om verantwoord gebruik te kunnen maken van AI-beeldgeneratietools is het cruciaal dat journalisten zich ervan vergewissen dat het trainingsmateriaal rechtmatig verkregen is.

*op basis van de interviews en het literatuuronderzoek

5.2.10 Vervolgonderzoek

Op basis van dit rapport formuleren we de volgende aanbevelingen voor vervolgonderzoek, om zo een steviger fundament te leggen voor het verantwoord en effectief inzetten van AI in de journalistieke beeldproductie.

Een diepgaander vergelijkend onderzoek naar de perceptie van AI-gegenereerde beelden onder verschillende doelgroepen zou waardevol zijn. Hoe verschillen de meningen tussen journalisten, het brede publiek en factcheckers over de betrouwbaarheid van AI-beelden? Daarnaast is een longitudinale studie interessant om te achterhalen hoe deze percepties zich in de loop van de tijd ontwikkelen, naarmate AI-toepassingen gangbaarder worden in de journalistiek.

Een tweede aanbeveling is een analyse van de impact van AI-beeldgebruik op het journalistieke vertrouwen. Onderzoek kan zich richten op de vraag of en hoe nieuwsconsumenten anders reageren op nieuwsartikelen met AI-beelden in vergelijking met traditionele fotografie, en in hoeverre dit invloed heeft op de geloofwaardigheid van nieuwsorganisaties.

Verder is er onderzoek nodig naar de effectiviteit van bestaande en toekomstige richtlijnen voor AI-gebruik in de journalistiek. Welke redacties hanteren al een beleid, en hoe wordt dit in de praktijk toegepast? Welke richtlijnen werken goed, en waar ontstaan er problemen of hiaten? Hierbij kan gekeken worden naar zelfregulering versus externe regulering, en de balans tussen ethische principes en journalistieke autonomie.

Ook zou het zinvol zijn om het juridisch kader rond AI-gegenereerd beeld verder te exploreren. Hoe evolueert de wetgeving rond auteursrechten, transparantieverplichtingen en aansprakelijkheid in verschillende rechtsgebieden? En hoe verhoudt die zich tot journalistieke vrijheden en de noodzaak van innovatieve beeldproductie?

Een bijkomende piste voor verder onderzoek is de ecologische impact van AI-gegenereerde beeldproductie. Het trainen en toepassen van AI-modellen vergt aanzienlijke rekenkracht en energieverbruik, wat bijdraagt aan de klimaatimpact van digitale media. Hoe verhoudt de energieconsumptie van AI-beeldgeneratie zich tot die van traditionele fotografie en beeldbewerking? Kunnen nieuwsredacties duurzamere keuzes maken bij het inzetten van AI in hun visuele content?

Ten slotte zou een praktijkgericht experiment met journalisten en beeldredacteuren kunnen uitwijzen in hoeverre AI-tools al daadwerkelijk geïntegreerd kunnen worden in het productieproces, en hoe redacties AI kunnen inzetten zonder in te boeten op transparantie en betrouwbaarheid. Wat zijn best practices voor het trainen van journalisten in verantwoord AI-gebruik? En hoe kunnen redacties AI-gegenereerde beelden verifiëren en annoteren op een manier die voor het publiek helder en vertrouwd aanvoelt?

Referenties

- Alenichev, A., Kingori, P., & Grietens, K. P. (2023). Reflections before the storm: the AI reproduction of biased imagery in global health visuals. *The Lancet Global Health*, 11(10), 1496-1498.
- Altay, S., & Gilardi, F. (2024). People are skeptical of headlines labeled as AI-generated, even if true or human-made, because they assume full AI automation. *PNAS Nexus*, 3(10), 403.
- Arends, S., & de Jong, J. (2023). The conflicting functions of image editing guidelines in journalism. *SocArchiv*. <https://doi.org/10.31235/osfio/bz79j>
- Ash, E., Durante, R., Grebenshchikova, M., & Schwarz, C. (2021). *Visual representation and stereotypes in news media* (SSRN Working Paper No. 3934858). Social Science Research Network. <https://dx.doi.org/10.2139/ssrn.3934858>
- Asmelash, S., Remmers, J., & Vang, T. (2018). *Mass Media: The Construction of Ethnic Stereotypes. Humanity in Action*, <http://www.humanityinaction.org/knowledgebase/557-mass-media-the-construction-of-ethnic-stereotypes>, Son Erisim Tarihi, 10.
- Balçık, T. (2019). *Moslims in Nederlandse kranten*. The Hague Peace Projects. [https://www.republiekallochtonie.nl/userfiles/files/2019-03-19-Rapport-Moslims-in-NL-kranten-THPP-NW-RA-1\(1\).pdf](https://www.republiekallochtonie.nl/userfiles/files/2019-03-19-Rapport-Moslims-in-NL-kranten-THPP-NW-RA-1(1).pdf)
- Bender, S. (2024). Generative-AI, the media industries, and the disappearance of human creative labour. *Media Practice and Education*, 1-18.
- Bianchi, F., Kalluri, P., Durmus, E., Ladhak, F., Cheng, M., Nozza, D., ... & Caliskan, A. (2023, June). Easily accessible text-to-image generation amplifies demographic stereotypes at large scale. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency* (pp. 1493-1504).
- Boyles, J. L., & Meisinger, J. (2020). Automation and adaptation: Reshaping journalistic labor in the newsroom library. *Convergence*, 26(1), 178-192.
- Broussard, M., Diakopoulos, N., Guzman, A. L., Abebe, R., Dupagne, M., & Chuan, C. H. (2019). Artificial intelligence and journalism. *Journalism & Mass Communication Quarterly*, 96(3), 673-695.
- Chauhan, A., Anand, T., Jauhari, T., Shah, A., Singh, R., Rajaram, A., & Vanga, R. (2024, February). Identifying race and gender bias in stable diffusion AI image generation. In *2024 IEEE 3rd International Conference on AI in Cybersecurity (ICAIC)* (pp. 1-6). IEEE.
- Clayton, A. (2023, 23 mei). Fake AI-generated image of explosion near Pentagon spreads on social media. *The Guardian*. <https://www.theguardian.com/technology/2023/may/22/pentagon-ai-generated-image-explosion>

- Cools, H., & Diakopoulos, N. (2023, 9 juli). Towards guidelines for guidelines on the use of generative AI in newsrooms. *Medium*. <https://generative-ai-newsroom.com/towards-guidelines-for-guidelines-on-the-use-of-generative-ai-in-newsrooms-55b0c-2c1d960>
- Dierickx, L., Sirén-Heikel, S., & Lindén, C. G. (2024). Outsourcing, augmenting, or complicating: The dynamics of AI in fact-checking practices in the Nordics. *Emerging Media*, 2(3), 449-473.
- Dignum, V., Régis, C., Bach, K., Buijsman, S., de Carvalho, A. P. L. F., Castellano, G., Dignum, F., Farries, E., Giannotti, F., Han, T. A., Helberger, N., Hellegren, I., Houben, G. J., Jahn, A., Joshi, S., Lamine Sarr, M., Lewis, D., Lind, A. S., Lugangira, N., ... Bourguine de Meder, Y. (2025). *Roadmap for AI policy research*. <https://www.researchgate.net/publication/388918982>
- Dodds, T., Vandendaele, A., Simon, F. M., Helberger, N., Resendez, V., & Yeung, W. N. (2025). Knowledge silos as a barrier to responsible AI practices in journalism? Exploratory evidence from four Dutch news organisations. *Journalism Studies*, 1-19.
- Fu, Y., & Stasko, J. (2023). More than data stories: Broadening the role of visualization in contemporary journalism. *IEEE Transactions on Visualization and Computer Graphics*, 1-20.
- Georgiou, M., & Zaborowski, R. (2017). *Media coverage of the “refugee crisis”: A cross-European perspective*. Council of Europe Report: DG1(2017)03. <https://edoc.coe.int/en/refugees/7367-media-coverage-of-the-refugee-crisis-a-cross-europeanperspective.html>
- Grabe, M. E., & Bucy, E. P. (2009). *Image bite politics: News and the visual framing of elections*. Oxford University Press.
- Grittmann, E. (2019). Methoden der Medienbildanalyse in der visuellen Kommunikationsforschung: Ein Überblick. In M. Löffelholz, L. Rothenberger, & A. Scheu (Eds.), *Handbuch visuelle Kommunikationsforschung* (pp. 527-546). Springer VS.
- Gutiérrez-Caneda, B., Lindén, C. G., & Vázquez-Herrero, J. (2024). Ethics and journalistic challenges in the age of artificial intelligence: Talking with professionals and experts. *Frontiers in Communication*, 9, 1465178.
- Gynnild, A. (2019). Visual journalism. In T. P. Vos, & F. Hanusch (Eds.), *The international encyclopedia of journalism studies* (Vol. 3, P-Y, pp. 1630–1637). John Wiley & Sons.
- Hausken, L. (2024). Photorealism versus photography: AI-generated depiction in the age of visual disinformation. *Journal of Aesthetics & Culture*, 16(1), 2340787.
- Hoover Green, A., & Cohen, D. K. (2021). Centering human subjects: The ethics of “desk research” on political violence. *Journal of Global Security Studies*, 6(2), ogaa029.

- Jaspers, M. W., Steen, T., Van Den Bos, C., & Geenen, M. (2004). The think aloud method: A guide to user interface design. *International Journal of Medical Informatics*, 73(11-12), 781-795.
- Jones, B, Luger, E & Jones, R 2023, *Generative AI & journalism: A rapid risk-based review*. Edinburgh Futures Institute.
- Khan, S. A., Sheikhi, G., Opdahl, A. L., Rabbi, F., Stoppel, S., Trattner, C., & Dang-Nguyen, D. T. (2023). Visual user-generated content verification in journalism: an overview. *IEEE Access*, 11, 6748-6769.
- Koltermann, F. (2019). *Fotografie und Konflikt: Rezensionen und Kritiken*. BoD–Books on Demand.
- Komatsu, T., Gutierrez Lopez, M., Makri, S., Porlezza, C., Cooper, G., MacFarlane, A., & Missaoui, S. (2020, October). AI should embody our values: Investigating journalistic values to inform AI technology design. In *Proceedings of the 11th Nordic Conference on Human-Computer Interaction: Shaping Experiences, Shaping Society* (pp. 1-13).
- Lehmuskallio, A., Häkkinen, J., & Seppänen, J. (2019). Photorealistic computer-generated images are difficult to distinguish from digital photographs: A case study with professional photographers and photo-editors. *Visual Communication*, 18(4), 427-451.
- Li, J., Cao, H., Lin, L., Hou, Y., Zhu, R., & El Ali, A. (2024, May). User experience design professionals' perceptions of generative artificial intelligence. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems* (pp. 1-18).
- Luccioni, S., Jernite, Y., & Strubell, E. (2024, June). Power hungry processing: Watts driving the cost of ai deployment?. In *Proceedings of the 2024 ACM conference on fairness, accountability, and transparency* (pp. 85-99).
- Maslej, N., Fattorini, L., Perrault, R., Parli, V., Reuel, A., Brynjolfsson, E., Etchemendy, J., Ligett, K., Lyons, T., Manyika, J., Niebles, J. C., Shoham, Y., Wald, R., & Clark, J. (2024). *The AI Index 2024 Annual Report*. AI Index Steering Committee, Institute for Human-Centered AI, Stanford University. <https://aiindex.stanford.edu/report/>
- Matich, P., Thomson, T. J., & Thomas, R. J. (2025). Old threats, new name? Generative AI and visual journalism. *Journalism Practice*, 1-20.
- Metz, C., & Schmidt, G. (2023, 29 maart). Elon Musk and others call for pause on A.I., citing 'profound risks to society'. *The New York Times*. <https://www.nytimes.com/2023/03/29/technology/ai-artificial-intelligence-musk-risks.html>
- Messaris, P., & Abraham, L. (2001). The role of images in framing news stories. In *Framing public life* (pp. 231-242). Routledge.
- Michel, T. (2023, 25 augustus). Hoe een AI-gegenereerde foto in het NOS Journaal terecht kwam. NOS. <https://over.nos.nl/nieuws/hoe-een-ai-gegenereerde-foto-in-het-nos-journaal-terecht-kwam/>

- Mortensen, T. M., McDermott, B. P., & Ejaz, K. (2023). Measuring photo credibility in journalistic contexts: Scale development and application to staff and stock photography. *Journalism Practice*, 17(6), 1158-1177.
- Newton, J. H. (1998). The burden of visual truth: The role of photojournalism in mediating reality. *Visual Communication Quarterly*, 5(4), 4-9.
- Nightingale, S. J., & Farid, H. (2022). AI-synthesized faces are indistinguishable from real faces and more trustworthy. *Proceedings of the National Academy of Sciences*, 119(8), e2120481119.
- Nishal, S., & Diakopoulos, N. (2024). Envisioning the applications and implications of generative AI for news media. *arXiv preprint*, arXiv:2402.18835.
- Noain-Sánchez, A. (2022). Addressing the impact of artificial intelligence on journalism: The perception of experts, journalists and academics. *Communication and Society*, 35(3), 105-121.
- Petca, A. R., Bivolaru, E., & Graf, T. A. (2012). Gender stereotypes in the Olympic Games media? A cross-cultural panel study of online visuals from Brazil, Germany and the United States. *Sport in Society*, 16(5), 611-630.
- Pruchová Hruzová, A. (2021). What is the image of refugees in Central European media? *European Journal of Cultural Studies*, 24(1), 240-258.
- Prüst, M. (2025). Fotografie en de Cultuurmonitor. Marc Prüst Consultancy. <https://www.cultuurmonitor.nl/sector/fotografie>
- Schröter, J. (2023). The AI image, the dream, and the statistical unconscious. *IMAGE. Zeitschrift für Interdisziplinäre Bildwissenschaft*, 19(1), 112-120.
- Strauss, A., & Corbin, J. (1990). *Basics of qualitative research (Vol. 15)*. Sage.
- Taha, S., Aissani, R., & Abdallah, R. (2024, September). AI applications for news. In 2024 *International Conference on Intelligent Computing, Communication, Networking and Services (ICCNS)* (pp. 220-227). IEEE.
- Thomson, T. J. (2019). To see and be seen: *The environments, interactions and identities behind news images*. Rowman & Littlefield.
- Thomson, T. J., Thomas, R. J., & Matich, P. (2024). Generative visual AI in news organizations: Challenges, opportunities, perceptions, and policies. *Digital Journalism*, 1-22.
- Thomson, T. J., Thomas, R., Riedlinger, M., & Matich, P. (2025). *Generative AI & journalism: Content, journalistic perceptions, and audience experiences*. APO.
- Weber, W., & Rall, H. (2013). "We are journalists": Production practices, attitudes and a case study of the *New York Times* newsroom. *Interaktive Infografiken*, 161-172.

Wu, S. (2024). Journalists as individual users of artificial intelligence: Examining journalists' "value-motivated use" of ChatGPT and other AI tools within and without the newsroom. *Journalism*, 0(0), 1-19.

Zelizer, B. (2010). *About to die: How news images move the public*. Oxford University Press.

Bijlage 1 - Overzicht interviews experts

Functie	#	Organisatie
Onderzoeker	01	Universiteit van Amsterdam
Onderzoeker	02	Hogeschool van Amsterdam
Onderzoeker	03	Hogeschool van Amsterdam
Onderzoeker	04	International Centre of Photography School (New York City)
Fotograaf, auteur, 'fotodetective'	05	De Volkskrant
Onderzoeker	06	Artevelde Hogeschool
Hoofdredacteur	07	Omroep Brabant
Presentator (AI)	08	Omroep Brabant
Onderzoeker	09	Universiteit Gent
Onderzoeker	10	Universiteit Gent
Onderzoeker	11	Queensland University of Technology (Brisbane)
Onderzoeker & datajournalist	12	AlghoritmWatch
Coördinator desinformatie en factchecking	13	VRT Nieuws
Onderzoeker, bestuurslid	14	JournalismLab, NVJ
Onderzoeker	15	Rathenau Instituut
Onderzoeker, curator, strategist	16	Universiteit van Groningen, DuPho
Directeur, algemeen secretaris	17	NVJ
Jurist	18	DuPho

Bijlage 2 - Overzicht interviews praktijkmensen

Functie	#	N/R	Medium
Fotograaf	01	Nationaal	Getty Images
Fotograaf	02	Regionaal	Leidsch Dagblad, Mare
Fotograaf	03	Nationaal	NRC, ANP
Fotograaf	04	Nationaal	Volkskrant
Beeldredacteur	05	Nationaal	NOS
Fotograaf	06	Nationaal	Freelance
Beeldredacteur	07	Nationaal	De Correspondent
Grafisch vormgever	08	Nationaal	De Morgen
Grafisch vormgever	09	Nationaal	Freelance - Cretive Coder
Beeldredacteur	10	Nationaal	ANP
Grafisch vormgever	11	Nationaal	Beeld & Geluid
Beeldredacteur	12	Nationaal	Beeld & Geluid
Beeldredacteur	13	Nationaal	AD
Fotograaf	14	Nationaal	Freelance
Beeldredacteur	15	Regionaal	AD Regionaal
Grafisch vormgever	16	Regionaal	AD Regionaal
Beeldredacteur	17	Nationaal	ANP
Fotograaf	18	Nationaal	Volkskrant
Beeldredacteur	19	Nationaal	Financieele Dagblad
Fotograaf	20	Nationaal	Freelance
Beeldredacteur	21	Regionaal	AD
Grafisch vormgever	22	Nationaal	AD
Grafisch vormgever	23	Nationaal	NU.nl
Fotograaf	24	Nationaal	Freelance
Grafisch vormgever	25	Nationaal	Volkskrant
Fotograaf	26	Regionaal	Freelance
Ombudsman	27	Nationaal	Trouw
Ombudsman	28	Nationaal	Raad voor Journalistiek
Beeldredacteur	29	Nationaal	Buitenhof
Ombudsman	30	Regionaal	De Limburger
Ombudsman	31	Nationaal	NPO
Beeldredacteur	32	Nationaal	Het Parool
Fotograaf	33	Regionaal	NH Dagblad
Grafisch vormgever	34	Nationaal	EO
Fotograaf	35	Nationaal	NOS
Grafisch vormgever	36	Nationaal	KRO-NCRV
Fotograaf	37	Nationaal	NRC
Fotograaf	38	Nationaal	Volkskrant
Beeldredacteur	39	Nationaal	Trouw
Fotograaf	40	Nationaal	BNNVARA

Bijlage 3 - Overzicht AI-tools en hun functionaliteiten

Dit incomplete overzicht van publiek toegankelijke AI-beeldsoftware is bedoeld om een indruk te geven van dit zeer snel ontwikkelende terrein en de journalistieke relevantie ervan. De lijst hieronder bevat slechts een deel van de AI-tools die in maart 2025 online beschikbaar zijn.

Soort	Naam	Maker/ eigenaar	Gerelateerde tools	Functionaliteit	Open source (voor zover online te vinden)	Prijs	Website	Opmerkingen/gebruik door journalistieke media
Beeldgeneratie modellen die zich profileren als ethisch verantwoord	Adobe Firefly	Adobe	NVIDIA Picasso, Photoshop, Lightroom	Text-to-image + text-to-video (beta) beeldgeneratie via webapplicatie	Nee	Gratis, en een custom enterprise- licentie voor bedrijven	https://firefly.adobe.com/ en https://www.adobe.com/products/firefly/features/ai-video-generator.html (video)	Adobe publiceerde op zijn website een overzicht met de ethische principes die het bedrijf volgde bij de ontwikkeling van de beeldgenerator: https://www.adobe.com/ai/overview/ethics.html . Het bedrijf spreekt over videogenerator als "the first commercially safe video generation model". Het text-to-image model is volgens Adobe getraind op Flickr-beelden in het publieke domein en Adobe stock images. Maar niet iedere maker van beelden uit Adobe Stock is het daarmee eens, zij hebben namelijk geen toestemming gegeven voor gebruik van hun beelden voor AI-training: https://www.creativeblog.com/news/adobe-copyright-ai?ref=tess.ghost.io
	Mitsua Diffusion One	Virtual Youtuber Elan Mitsua		Kunstzinnige beelden getraind op beelden uit publiek domein + opt in-beelden. Niet fotorealistisch	Ja	Gratis, en enkel voor niet-commercieel gebruik	https://elanmitsua.com/en/data/	Mitsua is een van de weinige AI-bedrijven die de trainingsdata van het model publiek toegankelijk op zijn website heeft staan
	Makery.ai	vAlsual		Fotorealistische beelden van en AI-gegenereerde personen, en een omgeving om eigen AI-beeldgeneratiemodellen op te zetten	Nee	Gratis met watermerk en beperkt aantal beelden, \$9,99/maand voor beelden zonder watermerk	https://vaisual.com/about/	De makers geven op hun website zelf aan: "We have built the largest legally clean dataset marketplace for the AI industry and now offer custom dataset services to global innovators of AI"
	Tess	Tess.design	Stable Diffusion	Tool om met expliciete toestemming van artiesten d.m.v. een beperkt aantal referentiebeelden een AI-beeld te creëren. Training gebeurt in een model van Stable Diffusion dat van de grond is opgebouwd	Nee	Vanaf \$20/maand voor niet-commercieel gebruik, en \$40/maand voor commercieel gebruik	https://www.tess.design/	Volgens de eigen website "The first properly-licensed image generator that enables artists to own their style." Maar de tool werkt ook op basis van Stable Diffusion, dat erom bekend staat getraind te zijn met problematisch bronmateriaal. Niet te verwarren met 'Tess AI' van het bedrijf Pareto, een textgenerator
	Shutterstock AI Image Generator	Shutterstock	Ontwikkeld 'in collaboration with OpenAI, Meta, and LG AI Research'	Text-to-image beeldgeneratie	Nee	€6/maand voor 100 beeldgeneraties/maand voor commercieel gebruik	https://www.artificialintelligence-news.com/news/shutterstock-launches-ai-image-generator-with-ethical-focus/	Shutterstock geeft aan dat zijn AI-beeldgenerator wordt getraind met bronnen die 'de diversiteit van de wereld waarin we leven vertegenwoordigen'. Het bedrijf zegt dat het de bijdrage van menselijke artiesten vergoedt door hen royalty's te betalen
	Bria AI	Bria		Text-to-image beeldgeneratie + beeldbewerking	Sommige modellen, zoals het background-removal model zijn open source	Gratis tot 1000 generaties of bewerkingen per maand, daarna \$0,08 per bewerking	https://bria.ai/	Bria is volgens de eigen website getraind op "100% licensed data"
Andere beeldgeneratie modellen	Midjourney	Midjourney, Inc.		Text-to-image beeldgeneratie, toegankelijk via Discord-server	Nee	Gratis variant + betaalde variant tussen 10 en 120 dollar/maand (sneller en onbeperkt aantal beelden)	https://www.midjourney.com/home	Met Midjourney is het ook mogelijk om beelden van bestaande personen te creëren. Zie bijvoorbeeld dit 'prompt book' https://enchanting-trader-463.notion.site/Prompt-Database-0745ecf6b3a646c095cd0261957c97a4 Wordt door verschillende journalistieke media-organisaties gebruikt, waaronder The Economist (https://www.economist.com/news/2022/06/11/how-a-computer-designed-this-weeks-cover), Corriere della sera (https://web.archive.org/web/20230430161801/https://www.corriere.it/la-lettura/22_agosto_26/su-la-lettura-highsmith-inedita-citta-europee-trasformazione-af21c38-2522-11ed-9c55-2dc9b0266fc8.shtml), The Atlantic (https://www.theatlantic.com/newsletters/archive/2022/08/where-does-alex-jones-go-from-here/676843/) en Noorse publieke omroep NRK (per ongeluk: https://www.journalisten.no/nrk-brukte-ai-illustrasjon-ved-en-feil-i-omtaalen-av-strompriser/538674)
	Dall-E 3	Open AI	Microsoft Copilot/Bing image generator	Text-to-image beeldgeneratie	Nee	Inbegrepen bij ChatGPT Plus (€20/maand)	https://openai.com/index/dall-e-3/	Neigt naar minder fotorealistisch beeld dan bijvoorbeeld Midjourney. The Irish Times gebruikte per ongeluk beelden van deze AI-beeldgenerator in een artikel over het gebruik van zelfbruiningscrème. (https://www.theguardian.com/media/2023/may/14/irish-times-apologises-for-hoax-ai-article-about-womens-use-of-fake-tan)
	Stable diffusion	Stability AI		Text-to-image beeldgeneratie via hun programma Dreamstudio	Ja	Gratis voor bedrijven met tot 1 miljoen euro omzet. Anders is er een Enterprise license nodig, met individuele prijszetting	https://stability.ai/	Is getraind met de LAION-5B-database en weinig handmatige moderatie van trainingsdata, waardoor privé-informatie van personen uit trainingsbeelden en schadelijke beelden kunnen voorkomen in de gegenereerde beelden, net als gewelddadige beelden en seksueel expliciete beelden. CEO Emad Mostaque geeft aan: "[it is] peoples' responsibility as to whether they are ethical, moral, and legal in how they operate this technology" Getty Images spande al een rechtszaak aan in verband met copyrightscheiding. Die vindt in 2025 plaats.
	Grok Image Generator	X (xAI)	FLUX AI	Text-to-image beeldgeneratie		Gratis		Werd gebruikt door onder andere het NOS Jeugdjournaal in een experimentele animatievideo: (https://medium.com/nos-digital/animating-the-news-with-ai-2c4e56e8a274) en satirisch collectief United Unknown (https://reutersinstitute.politics.ox.ac.uk/news/how-ai-generated-disinformation-might-impact-years-elections-and-how-journalists-should-report) De nieuwe variant van GROK was ten tijde van dit onderzoek nog maar net uitgekomen, dus over journalistiek gebruik van beelden die ermee gemaakt zijn was nog weinig bekend

Microsoft Designer Image Creator	Microsoft		Text-to-image beeldgeneratie via hun programma Dreamstudio	Nee	Gratis beperkt aantal afbeeldingen of Copilot pro abonnement voor 22 euro per maand	https://designer.microsoft.com/image-creator	
ImageFX	Google/Alphabet Inc.	Imagen	Text-to-image beeldgeneratie	Nee		https://deepmind.google/technologies/imagen-3/	Werd getraind met onder meer de LAION-5B database, waarin veel afbeeldingen zonder toestemming werden opgenomen
Sora	Open AI	ChatGPT, Dall-E	Text-to-video-generator	Nee	Nog niet publiek toegankelijk	https://openai.com/index/sora/	Dit model lijkt goede video's te kunnen genereren, maar is nog in de testfase
Meta image generator	Meta		Text-to-image beeldgeneratie, onderdeel van Meta AI	Nee	Gratis	https://www.meta.ai/	Getraind op de publiek zichtbare afbeeldingen van Facebook. Niet beschikbaar in Europa wegens 'onvoorspelbare regulering' https://www.theguardian.com/technology/article/2024/jul/18/meta-release-advanced-ai-multimodal-llama-model-eu-facebook-owner Ethische richtlijnen staan hier: https://about.fb.com/news/2024/02/labeling-ai-generated-images-on-facebook-instagram-and-threads/?utm_source=www.therundown.ai&utm_medium=newsletter&utm_campaign=ai-deciphers-2-000-year-old-ancient-scroll
Canva AI	Canva, OpenAI	Dall-E, Imagen	Text-to-image beeldgeneratie, AI beeldbewerkingen zoals achtergrond verwijderen	Nee	Gratis	https://www.canva.com/ai-image-generator/	
Leonardo.ai	Leonardo Interactive Pty Ltd				Gratis	https://leonardo.ai/	
Craiyon	Craiyon LLC.	Dall-E Mini	Text-to-image beeldgeneratie	Nee	Gratis	https://www.craiyon.com/	
FLUX AI	Black Forest Labs	Freepik AI, DZine AI	Image-to-image en text-to-image beeldgeneratie	Ja	Gratis	https://www.dzine.ai/	Model dat eerder werd gebruikt door Grok (xAI), en nog steeds door Freepik en DZine. Het stable diffusion-onderzoeksteam verhuisde voor een groot deel naar dit bedrijf nadat Stable Diffusion in moeilijkheden kwam
Freepik AI	EQT	FLUX AI, DZine AI	Beeldbank met image-to-image en text-to-image beeldgeneratie	Ja	Gratis	https://www.freepik.com/ai/	Onderdeel van stock beeldbank Freepik, gebaseerd op Flux AI. Gebruikers van beeldbank Freepik klagen al dat er veel te veel AI-beelden op de site staan (model collapse): https://www.reddit.com/r/graphic_design/comments/138hk1b/whats_up_with_freepik_and_the_ai_generated_images/
Dzine AI	Dzine	FLUX AI, Freepik AI	AI beeldbewerkingssoftware met image-to-image en text-to-image beeldgeneratie	Ja	Gratis	https://www.dzine.ai/	Gebaseerd op FLUX AI
Aligndraw	Elman Mansimov		Text-to-image beeldgeneratie (32x32 pixels)	Niet bekend	Niet publiek beschikbaar	https://aligndraw.fellowship.xyz/	Wordt algemeen beschouwd als het eerste text-to-image beeldgeneratiemodel zoals we het kennen. Gelanceerd in 2015
Artbreeder	Morphogen	StyleGAN and BigGAN, thisfacedoesntotext	Text-to-image beeldgeneratie	Een kleinere versie van het model is open source beschikbaar	Gratis preview en abonnementen vanaf \$7,49/maand	https://www.artbreeder.com en https://morphogen.io/	AI een zeer vroeg beschikbaar model (2018), dat werkt met Generative Adversarial Networks
Visual Electric	Visual Electric Company		Text-to-image beeldgeneratie	Nee	Gratis + abonnementen van €96/maand	https://visualelectric.com/	Zeer fotorealistisch, en met een focus op marketingtoepassingen
Switchlight	Beeble, Inc.		Image-to-image beeldbewerking, waarmee op basis van geavanceerde beeldanalyse de belichting van foto's kan worden aangepast aan een andere omgeving of lichtsituatie	Niet bekend	Gratis (1024x1024 pixels) en Pro voor 12,99/maand (2048x2048 pixels)	https://www.switchlightbeeble.ai/	Interessant omdat het een mix wordt van 'echte foto's' en gegenereerde belichting, van erg indrukwekkende kwaliteit
Red Panda AI	Red Panda AI			Nee, maar wel source available	Gratis	https://redpandasai.com/pricing	
Recraft AI	Recraft, Inc.		Text-to-image beeldgeneratie + beeldbewerking, ontwerpen van logo's.	Niet bekend	Gratis + premium t/m 48 dollar/maand	https://www.recraft.ai/project/4c0b8b1f-6af5-4466-9c57-13bbc2081748	Zeer realistisch. Scoort het hoogst op het leaderboard van de Artificial Analysis Image Arena . "The image generator has an arena win rate of 72%, meaning users are more often than not picking the new model over its competitors"
Ideogram v2	Ideogram.ai	Ideogram	Text-to-image beeldgeneratie met leesbare tekst	Nee	Gratis + abonnementen tot 48\$/maand	https://ideogram.ai/login	Ideogram v2 scoort bijna net zo hoog als Midjourney op de AIA -ranking én heeft de vaardigheid om leesbare tekst te genereren
Amazon Titan	Amazon		Text-to-image beeldgeneratie	Nee	Een variabele prijs die wordt berekend in Amazon Bedrock-token	https://aws.amazon.com/bedrock/amazon-models/titan/	Onderdeel van AI-app bouwplatform Bedrock
D-ID	D-ID		Text-to-video beeldgeneratie, image to video beeldgeneratie	Nee	Gratis 5 minuten video, en abonnementen tot en met \$108/maand	https://www.d-id.com/about-us/	
Remini	Bending Spoons S.p.A.		Lage kwaliteit foto's én video's 'verbeteren': o.a. resolutie verhogen, ruis weghalen, gezicht glad maken	Nee	Gratis basisfunctionaliteit of Pro-abonnement voor 9,99/week	https://remini.ai/	
Lensgo			Text-to-video beeldgeneratie, video to video beeldgeneratie en 'style transfer'	Nee	Gratis voor video's tot 5 seconden, en betaalde abonnementen tussen de \$6,75 en \$45/maand, afhankelijk van de tijdsduur van de gegenereerde video's	https://lensgo.ai/	
Genmo Mochi 1	Genmo		Text to video beeldgeneratie	Ja	Gratis, maar beperkt aantal video's per dag + meer video's voor \$96 tot en met \$288 per jaar		

	RunwayML	Runway		Text to video beeldgeneratie	Nee	Gratis trial, betaalde abonnementen tussen de \$12 en \$76 per gebruiker per maand	https://app.runwayml.com/video-tools/	
	Luma Dream Machine	Lumalabs		Text-to-image + text-to-video beeldgeneratie via webapplicatie	Nee	Gratis genereren van afbeeldingen (720p), niet-commerciële videogeneratie voor \$6,99/maand. Commercieel gebruik voor \$20,99/maand	https://dream-machine.lumalabs.ai/	
	Magnific AI	Magnific		Beeld opschalen, text-to-image beeldgeneratie	Nee	\$39/maand	https://magnific.ai/?utm_source=aitoolreport&utm_medium=newsletter&utm_campaign=u-s-uk-16-others-sign-ai-agreement	
Beeldzoekprogramma's die werken met AI	Associated Press AI-powered search	Associated Press		Video- en foto-analyse. D.m.v. AI worden video's doorzoekbaar gemaakt op inhoud, waardoor je makkelijker een specifiek fragment kan vinden		Geen aparte kosten	https://newsroom.ap.org/editorial-photos-videos/home	Controversieel in verband met privacyprobleem. Je kan er vreemde mensen op straat mee identificeren. https://www.nytimes.com/2023/10/23/technology/pimeyes-blocks-searches-childrens-faces.html , aangeklaagd in de VS, VK en Duitsland
	PimEyes	EMEARobotics		Reverse image search, gespecialiseerd in gezichtsherkenning		35,99 t/m 351,99 euro per maand	https://pimeyes.com/en	
	Lenso AI	Lenso AI S.A.		Reverse image search		7,51 euro t/m 181 euro per maand	https://lenso.ai/en	
	Lummi			Beeldbank met door AI gegenereerde stockfoto's, die ook nog eens kunnen worden ge 'reframed'		Gratis of \$8 per maand voor hoge resolutie	https://www.lummi.ai/	
Captiongenerators	Alt text hero	Center for Cooperative Media, Montclair State University	ChatGPT	Image-to-text: genereren van Alt-text of captions op basis van een beeld. Werkt binnen ChatGPT	Nee	Gratis	https://chatgpt.com/g/g-9qzzER8zo-alt-text-hero	
	Alt-text-helper	Good good good	ChatGPT	Image-to-text: genereren van Alt-text of captions op basis van een beeld. Werkt binnen ChatGPT	Nee	Gratis	https://chatgpt.com/g/g-vluyxEl3F-alt-text-helper	
AI-detectietools	Adobe Content Authenticity	Adobe	Onderdeel van de coalitie C2PA	Chrome browser-extensie die afbeeldingen een betrouwbaarheidsscore geeft.	N.v.t.	Gratis	https://chromewebstore.google.com/detail/adobe-content-authenticity/dmfbmenkapmaoldfgacgkqaoibikimel	Content Credentials are a new kind of tamper-evident metadata that serves as a "nutrition label" for digital content. Content Credentials can help you learn more about the content you consume online and provide you with information that can help you decide how to interpret it, whether you trust it, and how to stay more informed about it
	Forensically			Webapplicatie om onregelmatigheden in foto's mee te detecteren	N.v.t.	Gratis	https://29a.ch/photo-forensics/#clone-detection	
	Brandwell.ai			Geeft een 'waarschijnlijkheidsscore' over of een beeld met AI gemaakt is	N.v.t.	Gratis	https://brandwell.ai/ai-image-detector/?fp=ddiv	
Andere	Nightshade	Shawn Shan, Wenxin Ding, Josephine Passananti, Stanley Wu, Haitao Zheng, Ben Y. Zhao		Data poisoning: Sabotage van text-to-image tools door middel van 'vergiftigde beelden' met foute metadata en verwarrende pixels	N.v.t.	N.v.t.	https://arxiv.org/abs/2310.13828 en https://www.technologyreview.com/2023/10/23/1082189/data-poisoning-artists-fight-generative-ai/	How do these models compare? Are the open source alternatives on par with their proprietary counterparts? The Artificial Analysis Text to Image Leaderboard aims to answer these questions with human preference based rankings. The ELO score is informed by over 45,000 human image preferences collected in Artificial Analysis Image Arena. The leaderboard features the leading open-source and proprietary image models: the latest versions of Midjourney OpenAI's DALL-E, Stable Diffusion, Playground and more"
	Artificial Analysis Image Arena			Ranking van text-to-image modellen	N.v.t.	N.v.t.	https://huggingface.co/spaces/ArtificialAnalysis/Text-to-Image-Leaderboard en https://huggingface.co/spaces/ArtificialAnalysis/Text-to-Image-Leaderboard	
	Fjorney			Bot om automatisch een wachtrij aan Midjourney prompts uit te laten voeren. O.a. via Google Chrome-extensie	N.v.t.	Gratis (20 prompts per dag) en 2,89\$/maand (onbeperkt aantal prompts)	https://fjorney.com	
	FaceCheck			AI gezichtsherkenning	N.v.t.	Gratis gezichtsherkenning, maar om de identiteit van de herkende persoon te weten te komen moet men betalen, vanaf \$19	FaceCheck.id	
	Picarta			AI-geolocatie op basis van een beeld	N.v.t.	Gratis tot en met 3 zoekopdrachten per dag.	https://picarta.ai/	
	JournalismAI			Netwerk gesponsord door Google News, dat zich richt op 'verantwoord' gebruik van AI-tools door journalisten	N.v.t.	N.v.t.	https://www.journalism.ai/info/	

AI in Beeld

Fairly Trained Certified-keurmerk			Keurmerk voor 'eerlijk getrainde' AI, maar vooral voor AI-muziek	N.v.t.	N.v.t.	https:// www.fairlytrained.org/	Gesponsord door o.a. Universal, Authors Guild, Artists Rights Society en Association of Independent Music Publishers
Midlibrary			Een bibliotheek voor stijlen van bekende artiesten in Midjourney	N.v.t.	Gratis	https:// midlibrary.io/ art-styles	

