

Innovatief, maar overtuigend?

De invloed van functionele en emotionele reclameappeals op merkbetrouwbaarheid bij AI-gegenereerde reclame onder Generatie Z en de modererende rol van merkbekendheid



Juul van de Kamp

Communicatiewetenschap

Master Thesis

Innovatief, maar overtuigend?

De invloed van functionele en emotionele reclameappeals op merkbetrouwbaarheid bij AI-gegenereerde reclame onder Generatie Z en de modererende rol van merkbekendheid

Juul van de Kamp

S1098323

juul.vandekamp@ru.nl

Datum: 22 juni 2026

Aantal woorden: 9870

Masterthesis Communicatiewetenschap

Master Communicatiewetenschap

Specialisatie Commerciële Communicatie

Radboud Universiteit Nijmegen

Begeleider: dr. P.E. Ketelaar

Voorwoord

Wat gebeurt er met vertrouwen in een merk wanneer een reclame niet langer volledig door mensen, maar met behulp van AI wordt gemaakt? Die vraag vormde voor mij het startpunt van mijn thesis en bleef gedurende het hele onderzoeksproces terugkomen. Niet alleen tijdens het lezen van wetenschappelijke artikelen of het uitvoeren van analyses, maar ook in alledaagse situaties waarin AI-gegenereerde reclame ineens overal zichtbaar bleek te zijn.

Mijn interesse in dit onderwerp is ontstaan doordat ik in het dagelijks leven steeds vaker AI-advertenties voorbij zag komen, zoals bij een AI-kerstcampagne van McDonald's die online kritisch werd ontvangen. Ik zag dat AI niet alleen werd gebruikt door grote merken, maar ook door kleinere organisaties. Zo kwam ik bij mijn eigen voetbalclub een AI-gegenereerde poster tegen voor een trainingsset die voor kerst in de aanbieding was. De poster was duidelijk snel en efficiënt gemaakt, maar aan de vervormde letters in het logo was ook duidelijk zichtbaar dat AI was gebruikt.

Juist dat voorbeeld maakte het onderwerp voor mij interessant. Ik begreep meteen goed waarom een kleine organisatie AI zou inzetten: het is efficiënt, snel beschikbaar en vrijwel kosteloos. Tegelijkertijd riep het bij mij de vraag op wat zo'n AI-advertentie doet met de betrouwbaarheid van het merk of de organisatie erachter. Niet omdat AI per definitie goed of slecht is, maar omdat de manier waarop AI wordt ingezet waarschijnlijk bepaalt hoe consumenten de boodschap beoordelen. AI wordt volop gebruikt en zal waarschijnlijk een steeds normaler onderdeel worden van commerciële communicatie. Daarom wilde ik onderzoeken op wat voor manier AI-gegenereerde reclame effectief kan worden toegepast.

Graag bedank ik mijn begeleider Paul Ketelaar voor zijn kritische blik, heldere feedback en inhoudelijke betrokkenheid tijdens dit proces. Ook wil ik mijn medestudenten uit de masterclass bedanken voor het meedenken en de feedback die zij mij hebben gegeven tijdens de masterclasses.

Met deze thesis sluit ik mijn master-opleiding af met een onderwerp dat denk ik actueler is dan ooit. AI verandert de manier waarop reclame wordt gemaakt, maar deze studie laat voor mij vooral zien dat technologische vernieuwing in reclame niet altijd positief wordt ontvangen door consumenten. Juist in een tijd waarin communicatie steeds slimmer en sneller wordt, blijft de vraag naar geloofwaardige communicatie misschien wel belangrijker dan ooit.

Juul van de Kamp

Samenvatting

Generatieve AI wordt steeds vaker ingezet binnen commerciële communicatie. Het is echter nog onduidelijk onder welke omstandigheden AI-gegenereerde reclame door jonge consumenten uit Generatie Z als betrouwbaar wordt ervaren. Deze studie onderzocht in hoeverre reclameappeal (functioneel versus emotioneel) invloed heeft op de perceptie van merkbetrouwbaarheid bij AI-gegenereerde reclame, en of merkbekendheid dit effect modereert. In een 2 (appeal: functioneel versus emotioneel x 2 (merkbekendheid: bekend versus onbekend merk) between-subjects-experiment werden respondenten uit Generatie Z (N = 143) willekeurig toegewezen aan condities met een AI-gegenereerde advertentie voor een huidverzorgingsproduct. Na het bekijken van de advertentie beantwoordden zij vragen over de waargenomen merkbetrouwbaarheid.

De resultaten laten zien dat functionele AI-gegenereerde advertenties tot een significant hogere perceptie van merkbetrouwbaarheid leidden dan emotionele AI-gegenereerde advertenties. De modererende rol van merkbekendheid werd niet ondersteund: het effect van het reclameappeal op merkbetrouwbaarheid was onafhankelijk van of het merk bekend was bij de respondenten. Deze bevindingen suggereren dat bij AI-gegenereerde reclame vooral de aansluiting tussen de technologische aard van AI en de inhoud van de boodschap relevant is.

Dit onderzoek biedt hiermee praktische inzichten voor merken die generatieve AI willen inzetten in reclame en draagt bij aan de literatuur over AI in advertising, merkbetrouwbaarheid en reclame-effectiviteit onder Generatie Z. De bevindingen suggereren dat de effectiviteit van AI-gegenereerde reclame niet alleen afhangt van de inzet van AI zelf, maar ook van de manier waarop de reclameboodschap inhoudelijk wordt vormgegeven.

Inhoudsopgave

Hoofdstuk 1 – Inleiding	6
Hoofdstuk 2 – Theoretisch kader	8
2.1 De relatie tussen reclameappeal en merkbetrouwbaarheid	8
2.2 De modererende invloed van merkbekendheid op de relatie tussen reclameappeal en merkbetrouwbaarheid.....	10
Hoofdstuk 3 – Methode van onderzoek	12
3.1 Onderzoeksdesign	12
3.2 Onderzoekspopulatie en steekproef	12
3.3 Onderzoeksmateriaal	13
3.4 Meetinstrument	15
3.5 Procedure.....	17
3.6 Data-analyse	18
H4 – Resultaten	20
4.1 Schaalconstructie Merkbetrouwbaarheid	20
4.2 Beschrijvende gegevens	20
4.3 Check op aannames	21
4.4 Covariaten	22
4.5 Hypothesetoetsing	22
4.6 Aanvullende analyses.....	24
H5 – Conclusie en discussie.....	26
5.1 Conclusie	26
5.2 Theoretische terugkoppeling van de bevindingen	27
5.3 Alternatieve theoretische verklaringen voor de onderzoeksresultaten	29
5.4 Methodologische beperkingen	30
5.5 Suggesties voor vervolgonderzoek	32
5.6 Wetenschappelijke implicaties	33
5.7 Praktische implicaties	34
Literatuurlijst	36
Bijlage A – Informatiebrief en toestemmingsformulier.....	40
Bijlage B – Vragenlijst.....	43

Hoofdstuk 1 – Inleiding

Generatieve AI is haast niet meer weg te denken uit de wereld van reclame. Waar reclame traditioneel werd ontwikkeld door mensen, wordt commerciële communicatie steeds vaker deels gemaakt met hulp van kunstmatige intelligentie (Have, 2026). AI kan worden gebruikt voor het genereren van beelden, teksten en zelfs complete campagnes (Weatherbed, 2026). Steeds meer bedrijven experimenteren met AI en zetten deze technologie actief in voor hun reclame-uitingen (Weatherbed, 2026). Hoewel deze recente ontwikkelingen merken nieuwe kansen biedt, roept het gebruik van AI in reclame ook vragen op over geloofwaardigheid, authenticiteit en vertrouwen. Juist om deze reden is het belangrijk om te onderzoeken onder welke omstandigheden AI-gegenereerde reclame bijdraagt aan merkbetrouwbaarheid.

Deze vraag is vooral relevant in het kader van de doelgroep Generatie Z. Generatie Z is de eerste leeftijdsgroep die is opgegroeid in een wereld die door technologie wordt gedomineerd, daarnaast zijn Generatie Z'ers de allereerste digital natives (Segura, 2024). Het komt daarom niet als verrassing dat veel Generatie Z'ers AI gebruiken. Uit onderzoek van (Harvard Business Impact Education, 2026) blijkt dat 74% van Generatie Z meer dan eens in de maand gebruikmaakt van AI-chatbots. Om deze reden denken adverteerders dat Generatie Z positief of neutraal staat tegenover het gebruik van generatieve AI in reclame (IAB, 2024). Dit blijkt echter niet vanzelfsprekend, slechts 48% van millennials en Generatie Z geeft aan 'enigszins positief' te zijn over AI in reclame (IAB, 2024). Volgens IAB (2024) komt deze kritische houding onder andere doordat AI-reclame als onecht, oninteressant of onethisch kan worden ervaren. Deze twijfels raken direct aan de authenticiteit, intentie en geloofwaardigheid van de boodschap, waardoor merkbetrouwbaarheid een relevante uitkomstvariabele is.

Voor bedrijven is het daarom cruciaal om te weten onder welke omstandigheden generatieve AI in reclame de waargenomen merkbetrouwbaarheid kan ondersteunen of juist ondermijnen. Een belangrijke strategische keuze hierbij is het type reclameappeal: een functioneel appeal, gericht op concrete productinformatie en rationele argumenten, of een emotioneel appeal, gericht op gevoelens en affectieve reacties. Hoewel onderzoek suggereert dat functionele appeals effectiever kunnen zijn bij AI-gegenereerde content (Song et al., 2024), is het onduidelijk of deze voorkeur standhoudt binnen AI-genereerde reclame onder Generatie Z en in hoeverre dit doorwerkt in de perceptie van merkbetrouwbaarheid.

Naast het type appeal kan de relatie tussen consument en merk van invloed zijn. Merkbekendheid, de mate waarin consumenten een merk herkennen en associaties hebben

kan mogelijk fungeren als een buffer voor twijfels over authenticiteit (Ha & Perks, 2005). Bij een bekend merk beschikken consumenten vaak al over bestaande merkassociaties, terwijl een onbekend merk vanwege het gebrek aan merkassociaties sterker afhankelijk kan zijn van de advertentie zelf. Het is echter onbekend of deze merkkracht de perceptie van AI-gegenereerde reclame bij Generatie Z beïnvloedt. Daarmee is het nog onduidelijk of merkbekendheid de invloed van reclameappeal op merkbetrouwbaarheid versterkt of juist verzwakt.

Opvallend is dat veel bestaand onderzoek AI-gegenereerde reclame nog altijd benadert vanuit een vergelijking met menselijke reclame. Generatieve AI wordt inmiddels steeds meer geïntegreerd in marketingcommunicatie (Henke, 2025) en AI ontwikkelt zich van experimenteel hulpmiddel tot een steeds normaler onderdeel van de branche. Waar voorgaande studies voornamelijk de fundamentele vergelijking tussen AI en menselijke makers hebben onderzocht, bevindt de marketingpraktijk zich inmiddels in een fase waarin de inzet van AI de standaard vormt. Hierdoor verschuift de wetenschappelijke relevantie van de vraag óf AI moet worden ingezet naar de vraag onder welke strategische condities deze technologie de merkbetrouwbaarheid optimaal ondersteunt. Vanuit dat perspectief is het relevant om niet alleen AI met menselijke reclame te vergelijken, maar juist verschillen binnen AI-gegenereerde reclame te onderzoeken. Dit onderzoek richt zich daarom op de vraag hoe reclameappeal de perceptie van merkbetrouwbaarheid beïnvloedt, en of merkbekendheid deze relatie versterkt of verzwakt.

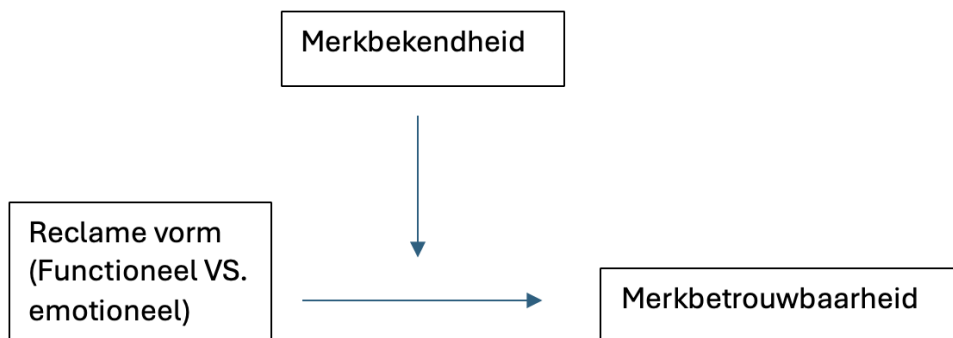
Dit onderzoek richt zich daarom op de volgende onderzoeksvraag:

In hoeverre beïnvloedt reclameappeal (functioneel versus emotioneel) de perceptie van merkbetrouwbaarheid bij AI-gegenereerde reclame onder Generatie Z, en in hoeverre wordt dit effect gemodereerd door merkbekendheid?

Hoofdstuk 2 - Theoretisch kader

In dit hoofdstuk wordt het theoretisch kader van het onderzoek uiteengezet. Allereerst wordt ingegaan op het verwachte hoofdeffect: hoe verschillende reclameappeals, emotionele en functionele appeals, van invloed kunnen zijn op de beoordeling van merkbetrouwbaarheid bij AI-gegenereerde reclame. Daarna wordt ingegaan op de modererende rol van merkbekendheid als factor die de invloed van reclameappeals op merkbetrouwbaarheid mogelijk beïnvloedt. Op basis van theoretische inzichten en empirisch onderzoek worden de hypotheses van dit onderzoek geformuleerd.

Figuur 1
Conceptueel model



2.1 De relatie tussen reclameappel en merkbetrouwbaarheid

In onderzoek worden reclameboodschappen vaak opgedeeld in verschillende types, ook wel Reclameappeals genoemd. Een Reclameappel verwijst naar de manier waarop een reclame een consument probeert te overtuigen. Twee stromingen hierin zijn emotionele en functionele appeals. Functionele appeals richten zich voornamelijk op objectieve informatie over het product, zoals functionaliteit en praktische voordelen. Emotionele appeals daarentegen proberen gevoelens of waarden bij de consument op te roepen (Leonidou & Leonidou, 2009).

Wanneer de bron van de boodschap geen menselijke maker is, maar AI, kan de werking van functionele en emotionele appeals veranderen. Vanuit de Machine Heuristic (Sundar & Kim, 2019) worden algoritmes primair geassocieerd met logica en dataverwerking. Hierdoor kan een functioneel appeal beter aansluiten bij de perceptie die consumenten hebben van AI als bron van een boodschap. Wu en Wen (2021) tonen aan dat consumenten AI-creaties positiever beoordelen wanneer het creatieproces wordt gezien als rationeel en datagedreven. Dit mechanisme kan nader worden geduid middels de Match-up

Hypothesis, die stelt dat de effectiviteit van een boodschap toeneemt wanneer de kenmerken van de bron overeenkomen met de aard van de gecommuniceerde inhoud (Lee et al., 2022). Wanneer de rationele aard van een algoritmische bron wordt gekoppeld aan een functionele boodschap, ontstaat er een hoge mate van congruentie. Vanuit deze theoretische basis is te redeneren dat de congruentie tussen de bron en de boodschap belangrijk is voor de perceptie van merkbetrouwbaarheid, omdat de rationele aard van de boodschap beter past bij de rationele en datagedreven associaties die consumenten met AI kunnen hebben. Een emotionele boodschap sluit mogelijk minder goed aan bij de associaties die mensen met AI hebben.

Eerder onderzoek toonde aan dat congruentie tussen verschillende elementen van een advertentie een belangrijke rol kan spelen in hoe een boodschap wordt geëvalueerd (Lee et al., 2022). Deze studie naar celebrity endorsements toont aan dat wanneer er een duidelijke match bestaat tussen de bron van een boodschap en het product of de boodschap zelf, dit kan bijdragen aan een hogere perceptie van broncredibiliteit en positievere attitudes ten opzichte van zowel de advertentie als het merk (Lee et al., 2022). In de context van AI-gegenereerde reclame kan een vergelijkbaar mechanisme optreden. Aangezien AI in de literatuur vaak wordt geassocieerd met rationaliteit, data-analyse en objectiviteit (Sundar & Kim, 2019), kan een functioneel reclameappeal beter aansluiten bij de perceptie van AI als bron van de boodschap.

Een emotioneel appeal kan daarentegen juist als incongruent worden beoordeeld. Emotionele reclame doet namelijk een beroep op gevoelens, empathie, oprechtheid en menselijke betrokkenheid, terwijl een AI-bron geen emoties kan ervaren en geen intrinsieke intenties of morele verantwoordelijkheid bezit. Vanuit deze redenering volgt dat een emotionele boodschap kan botsen met het gebrek aan menselijke intentie bij een AI-bron. Hierdoor kan bij consumenten het gevoel ontstaan dat de emotionele boodschap niet authentiek is, maar slechts algoritmisch wordt nagebootst. In plaats van vertrouwen te versterken, kan een emotionele AI-gegenereerde advertentie daardoor worden ervaren als kunstmatig, instrumenteel of zelfs manipulatief. Dit kan schadelijk zijn voor de perceptie van merkbetrouwbaarheid, omdat vertrouwen in een merk mede afhankelijk is van de mate waarin de communicatie als oprecht en geloofwaardig wordt ervaren.

Song et al. (2024) tonen aan dat AI-gegenereerde advertenties effectiever kunnen zijn wanneer zij gebruikmaken van een functioneel in plaats van een emotioneel appeal. Zij verklaren dit doordat functionele appeals beter aansluiten bij de objectieve en functionele associaties die consumenten met AI hebben. Omdat deze bevindingen contextafhankelijk

kunnen zijn, blijft het relevant om te onderzoeken of dit patroon ook optreedt bij AI-gegenereerde reclame onder Generatie Z.

H1: Bij Generatie Z leidt blootstelling aan een AI-gegenereerde advertentie met een functioneel appeal tot een hogere perceptie van merkbetrouwbaarheid dan bij blootstelling aan een AI-gegenereerde advertentie met een emotioneel appeal

2.2 De modererende invloed van merkbekendheid op de relatie tussen reclameappeal en merkbetrouwbaarheid

De effectiviteit van een reclameappeal staat niet op zichzelf, maar wordt bepaald door de interactie tussen het type reclameappeal en de mate van merkbekendheid. Consumenten beoordelen een advertentie niet enkel op basis van inhoud, maar ook op basis van de afzender van de boodschap. De effectiviteit van een appeal is daardoor onlosmakelijk verbonden met de reeds bestaande associaties die consumenten hebben bij het merk.

Merkbekendheid verwijst naar de mate waarin een merk aanwezig is in het geheugen van consumenten en kan daardoor fungeren als een filter voor nieuwe informatie (Keller, 1993). Wanneer een merk sterk aanwezig is in het geheugen van consumenten, beschikken zij doorgaans al over complexere kennisstructuren en associaties met betrekking tot dat merk die dienen als filter voor nieuwe informatie. Een hogere mate van merkbekendheid kan bijdragen aan merkvertrouwen, omdat consumenten bij bekende merken vaker beschikken over eerdere ervaringen en kennis over het merk (Ha & Perks, 2005). Daardoor kan een bekend merk fungeren als vertrouwensbasis bij de beoordeling van AI-gegenereerde reclame. Dit impliceert dat de invloed van het type appeal op de merkbetrouwbaarheid afhankelijk is van de mate waarin de afzender reeds als een vertrouwd referentiepunt fungeert.

In lijn met de studie van Keller (1993), tonen Laroche et al. (1996) aan dat merkbekendheid leidt tot een grotere ‘confidence in brand evaluations’. Deze zekerheid fungeert als een psychologisch schild: consumenten leunen op eerdere positieve ervaringen, waardoor de ‘algoritmische herkomst’ van AI-gegenereerde content minder zwaar weegt in de betrouwbaarheidsperceptie. Hierdoor kan het type reclameappeal bij bekende merken minder bepalend worden voor de perceptie van merkbetrouwbaarheid bij bekende merken in vergelijking met onbekende merken.

Een psychologisch mechanisme dat de relatie tussen merkbekendheid en de evaluatie van AI-boodschappen verder kan verklaren, is het mere exposure-effect. Volgens Zajonc (1968) kan herhaalde blootstelling aan een stimulus tot positievere evaluaties leiden, zelfs wanneer een individu deze stimulus maar oppervlakkig verwerkt. Dit betekent dat mensen objecten of merken die zij vaker tegenkomen vaak positiever beoordelen dan objecten

waarmee zij minder vertrouwd zijn. Bij onbekende merken ontbreekt deze vertrouwensbasis, waardoor een emotionele AI-appeal sneller als manipulatief wordt ervaren. Een functioneel appeal kan beter aansluiten bij onbekende merken, omdat deze vooral rationele productvoordelen benadrukt en minder beroep doet op een bestaande emotionele band tussen consument en merk.

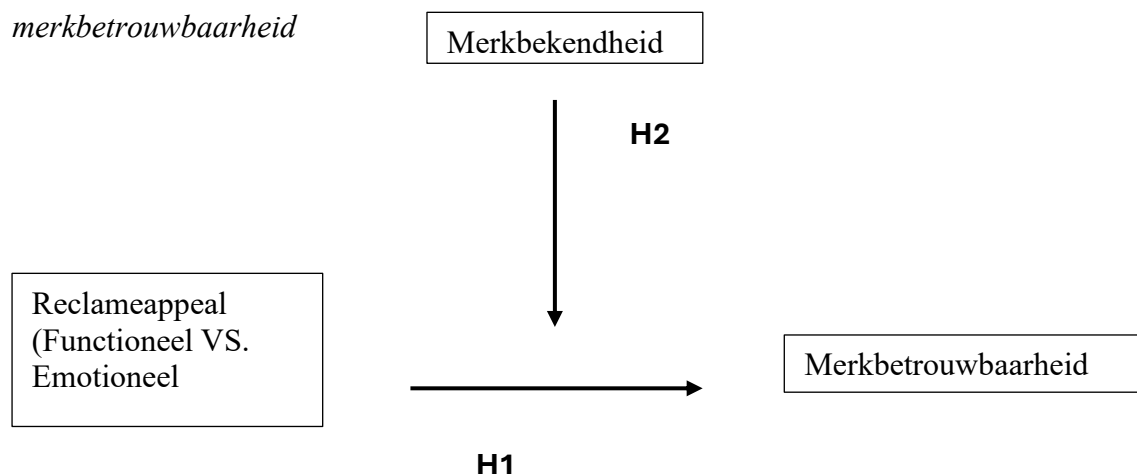
Het mechanisme waarbij merkbekendheid als buffer fungeert werkt anders bij een bekend merk. Generatie Z kan AI-content als minder authentiek of minder oprecht ervaren, waardoor de kans op scepsis ten opzichte van de reclameboodschap toeneemt. Wanneer een advertentie echter afkomstig is van een bekend merk, hoeven consumenten die boodschap niet volledig los van bestaande merkassociaties te beoordelen. Het merk kan dan fungeren als een vertrouwd ankerpunt, waardoor negatieve reacties mogelijk worden afgezwakt. De gevestigde reputatie compenseert hier mogelijk voor de vermeende gebrekkige authenticiteit van de AI-techniek. Hierdoor kan de noodzaak voor een strikt functionele insteek afnemen; het bekende merk is mogelijk in staat om met succes een emotionele AI-boodschap over te brengen zonder dat dit direct tot argwaan leidt.

Kortom, de verwachting is dat merkbekendheid het effect van reclameappeal op merkbetrouwbaarheid modereert: de positieve impact van een functionele (vs. emotionele) appeal is significant bij onbekende merken, maar vlagt af bij bekende merken.

Op basis van dit mechanisme wordt de volgende interactiehypothese geformuleerd:
H2: Merkbekendheid modereert het effect van het reclameappeal op de gepercipieerde merkbetrouwbaarheid bij Generatie Z, waarbij het positieve effect van een functioneel appeal sterker is bij merken met een lage merkbekendheid dan bij merken met een hoge merkbekendheid.

Figuur 2

Conceptueel model van de relatie tussen reclameappeal, merkbekendheid en merkbetrouwbaarheid



Hoofdstuk 3 – Methode van onderzoek

Dit hoofdstuk licht toe welke keuzes zijn gemaakt in het onderzoeksproces. Aan bod komen het onderzoeksdesign, de onderzoekspopulatie en steekproef, het onderzoeksmateriaal, de onderzoeksinstrumenten, de procedure en tot slot de data-analyse.

3.1 Onderzoeksdesign

Dit onderzoek maakt gebruik van een experiment. Een experiment maakt het mogelijk om causale relaties te onderzoeken binnen een gecontroleerde omgeving. Dit is essentieel voor het toetsen van de hypothesen op een manier die niet mogelijk is met bijvoorbeeld enkel survey-onderzoek of kwalitatieve methoden. Het experiment bestond uit een 2 (reclameappel: emotioneel versus functioneel) x 2 (merkbekendheid: fictief versus bekend) – between subjects design. Respondenten werden willekeurig toegewezen aan een van de vier verschillende condities. Deze opzet, waarbij de respondent maar 1 versie ziet van het stimulusmateriaal, zorgt ervoor dat het voor respondenten lastiger is om het doel van het onderzoek te achterhalen. Dit is goed voor de interne validiteit van het onderzoek (Orne, 1962).

De ethische goedkeuring is verkregen via de Ethische Commissie Sociale Wetenschappen van de Radboud Universiteit onder projectnummer: ECSW-LT-2026-3-29-58410.

3.2 Onderzoekspopulatie en steekproef

De onderzoekspopulatie die is gekozen voor dit onderzoek is Generatie Z. Uit eerder onderzoek blijkt dat Generatie Z zich bewust is van de kansen die generatieve AI biedt. Tegelijkertijd staan deze jongeren kritisch tegenover het gebruik van generatieve AI in reclame (Bachnik et al., 2025). Daarnaast is Generatie Z een generatie die is opgegroeid in een sterk gedigitaliseerde omgeving en daardoor vertrouwd is met technologie (Persada et al., 2019). Om die redenen is Generatie Z een relevante groep om te onderzoeken in de context van AI.

De steekproef is geworven via *convenience sampling*. Het onderzoek is verspreid via de persoonlijke netwerken van de onderzoeker, via verschillende sociale mediakanalen en het platform SONA. Deze manier van werven is gekozen vanwege de praktische haalbaarheid en de goede bereikbaarheid van de doelgroep via deze wervingsmethode. Daarnaast is gebruikgemaakt van *snowball sampling*; respondenten werden gevraagd de vragenlijst verder te verspreiden onder bekenden die eveneens tot de doelgroep behoren. Op deze manier kon het bereik van de vragenlijst worden vergroot. Hoewel deze steekproefmethoden niet leiden

tot een representatieve afspiegeling van Generatie Z of de Nederlandse bevolking als geheel, zijn zij passend voor dit experimentele onderzoek. De nadruk ligt namelijk niet op het schatten van populatiekenmerken, maar op het toetsen van causale mechanismen binnen een gecontroleerde experimentele setting.

Voorafgaand aan de dataverzameling is er een poweranalyse uitgevoerd om te bepalen hoeveel respondenten benodigd waren. Deze analyse is uitgevoerd om vast te stellen hoeveel deelnemers er nodig zijn om met voldoende statistische power een effect te kunnen detecteren. De analyse werd gedaan met G*Power versie 3.1.9.7 (Faul et al., 2009). Op basis van een F-test (ANOVA: Fixed effects, special, main effects and interactions) in G*Power is vastgesteld dat er minimaal 128 respondenten nodig zijn (32 per conditie). Hierbij is uitgegaan van een significantie van $\alpha = .05$, een gewenste power van .80 en een verwachte effectgrootte van $f = .25$. De verwachte effectgrootte is gebaseerd op eerder experimenteel onderzoek naar AI-gegenereerde reclame en advertising appeals. Song et al. (2024) vonden in meerdere 2×2 -experimenten significante interactie-effecten tussen creatortype en appeal, variërend van middelgrote tot grote effecten. Omdat dit onderzoek een andere context gebruikt, namelijk een vergelijking tussen verschillende vormen van AI-gegenereerde reclame in plaats van de vergelijking tussen AI en menselijke creators, is er gekozen voor een conservatieve middelgrote effectgrootte ($f = .25$) voor de poweranalyse.

In totaal namen 184 respondenten deel aan dit onderzoek. Met een screeningvraag naar leeftijd aan het begin van de vragenlijst is gecontroleerd of respondenten daadwerkelijk tot Generatie Z behoorden en op het moment van invullen minstens 16 jaar oud waren. Hoewel Generatie Z vaak wordt afgebakend als de generatie geboren tussen 1997 en 2012, is in dit onderzoek een ondergrens van 16 jaar gehanteerd vanwege ethische voorwaarden voor deelname. Alleen respondenten die tot de doelgroep behoorden en de vragenlijst volledig hadden ingevuld, zijn meegenomen in de analyse.

3.3 Onderzoeksmateriaal

Keuze product en merk

Voor dit onderzoek is gebruikgemaakt van stimulusmateriaal in de vorm van advertenties voor een huidverzorgingsproduct. Er is gekozen voor deze categorie, omdat het een productcategorie is die voor een brede doelgroep van toepassing is. Daarnaast leent deze productcategorie zich voor zowel een functioneel als een emotioneel reclameappeal. Een functioneel appeal richt zich hierbij op producteigenschappen, zoals in de advertentie de specifieke voordelen van gebruik en voor welke lichaamsdelen het geschikt is, terwijl een

emotioneel appeal zich bijvoorbeeld kan richten op gevoelens van zorg, zachtheid en vertrouwdeheid. In de gebruikte advertenties komt dit naar voren door de focus op de technische productvoordelen versus de affectieve beleving van het gebruik. Hierdoor was crème een geschikte productcategorie. Het gekozen product is een standaardcrème, een alledaags en relatief genderneutraal product. In tegenstelling tot sterk gesegmenteerde producten, zoals mannen- of vrouwenverzorging, maakt dit het mogelijk om een brede groep respondenten bloot te stellen aan hetzelfde stimulusmateriaal zonder dat het product slechts een deel van de steekproef aanspreekt. Respondenten zijn hierdoor voldoende in staat om een oordeel te vormen over het merk en de advertentie, zonder dat extreme voorkeuren of sterke betrokkenheid een rol spelen. Ook is merkbetrouwbaarheid een logische variabele binnen de productcategorie huidverzorging, omdat eerdere studies suggereren dat merkbetrouwbaarheid in deze markt een belangrijke rol speelt bij aankoopintenties (Sanny et al., 2020)

Het stimulusmateriaal

Het stimulusmateriaal bestond uit vier verschillende advertenties, passend bij het 2 x 2 between subjects-design. Hierbij werd reclameappeal gemanipuleerd door gebruik te maken van een meer emotionele en een meer functionele advertentie. Naast reclameappeal werd ook merkbekendheid gemanipuleerd. Hiervoor is gebruikgemaakt van een bekend merk, namelijk Nivea, en een fictief merk, namelijk Luxera. Er is gekozen voor Nivea, omdat dit een breed herkenbaar huidverzorgingsmerk is, waardoor de conditie met hoge merkbekendheid geloofwaardig en realistisch overkomt. Voor de conditie met lage merkbekendheid is een fictief merk ontwikkeld, zodat bestaande merkassociaties en eerdere ervaringen met het merk zoveel mogelijk worden uitgesloten. De keuze voor een fictief merk maakt het mogelijk om het effect van merkbekendheid beter te isoleren.

Bij het ontwikkelen van het stimulusmateriaal is geprobeerd om de reclame zo realistisch mogelijk te houden. Dit draagt bij aan de geloofwaardigheid en realistische relevantie van het onderzoek en vergroot de kans dat de bevindingen beter aansluiten bij daadwerkelijk consumentengedrag (Morales et al., 2017). Er is daarom de keuze gemaakt om twee bestaande advertenties te nemen van Nivea, een met een emotioneel appeal en een met een functioneel appeal. De advertenties zijn vervolgens aangepast voor Luxera. Bij het aanpassen van de originele advertentie is geprobeerd om visuele invloeden zoveel mogelijk constant te houden. Zo bleef binnen de condities de productcategorie, de algemene advertentie-opzet en de visuele focus op de crème verpakking zoveel mogelijk gelijk. Hierdoor kunnen verschillen in respons sterker worden toegeschreven aan de experimentele manipulaties van reclameappeal en merkbekendheid, in plaats van visuele verschillen.

Voordat respondenten het stimulusmateriaal zagen, lazen zij een korte introductietekst waarin de beoordelingscontext werd geschetst en werd vermeld dat de advertentie met behulp van AI was ontwikkeld. Alle respondenten ontvingen dezelfde introductietekst om de context van de beoordeling te standaardiseren. De volledige tekst is opgenomen in Bijlage B.

Tabel 1

Stimulusmateriaal advertentie

	Functioneel Appeal	Emotioneel Appeal
Merk: Nivea		
Merk: Luxera		

3.4 Meetinstrument

Dit onderzoek maakte gebruik van een bestaande, gevalideerde meetschaal. De items zijn vertaald naar het Nederlands en aangepast aan de context van dit experiment. Om de inhoudelijke betekenis van de oorspronkelijke items zoveel mogelijk te behouden, zijn de vertalingen gecontroleerd en besproken tijdens de masterclass. Op basis van deze feedback zijn formuleringen aangescherpt waar nodig. De stellingen voor merkbetrouwbaarheid zijn

gemeten op een 5-punts Likertschaal, variërend van 1 = helemaal mee oneens tot 5 = helemaal mee eens. Hogere scores wijzen telkens op een hogere mate van het gemeten construct.

Merkbetrouwbaarheid.

De afhankelijke variabele in dit onderzoek is merkbetrouwbaarheid. Deze variabele is gemeten met behulp van 10 items, ontleend aan de Brand Trust scale van Munuera-Aleman et al. (2003) en aangepast aan de context van dit experiment. Deze schaal is gekozen omdat deze specifiek is ontwikkeld voor het meten van vertrouwen in merken en toepasbaar is op verschillende productcategorieën. De oorspronkelijke schaal bevatte een enkel item dat betrekking had op langdurige gebruikservaring. Omdat respondenten in dit onderzoek slechts één advertentie te zien kregen en dus geen daadwerkelijke gebruikservaring met het merk of product opdeden, sloot dit item niet aan bij de experimentele context. Om deze reden zijn alleen de items geselecteerd die betrekking hebben op algemene percepties van betrouwbaarheid, eerlijkheid en welwillendheid.

Respondenten werd gevraagd in hoeverre zij het eens waren met uitspraken over het merk in de advertentie. Voorbeelditems zijn: “Dit merk voldoet aan mijn verwachtingen”, “Ik heb vertrouwen in dit merk”, “Dit merk komt eerlijk en oprecht over” en “Ik zou op dit merk kunnen vertrouwen als er een probleem ontstaat.” Hogere scores wijzen op een hogere mate van merkbetrouwbaarheid. De interne consistentie van de schaal bedroeg Cronbachs $\alpha = .839$.

Tabel 2*Meetinstrumentitems voor merkbetrouwbaarheid*

Item	Item stelling
A	Met merk [x] krijg ik wat ik zoek in een crème
B	[X] is een merk dat aan mijn verwachtingen voldoet.
C	Ik heb vertrouwen in merk [x].
D	[x] is een merk dat mij niet teleurstelt.
E	Merk [x] zou eerlijk en oprecht omgaan met mijn zorgen.
F	Merk [x] doet er alles aan om mij tevreden te stellen.
G	Ik zou op merk [x] kunnen vertrouwen om een probleem op te lossen.
H	Merk [x] vindt mijn tevredenheid belangrijk.
I	Merk [x] zou mij op de een of andere manier compenseren voor een probleem met de crème.
J	Merk [x] zou niet bereid zijn een probleem op te lossen dat ik met de crème zou kunnen hebben.

Sociodemografische variabelen. In dit onderzoek zijn sociodemografische variabelen, waaronder geslacht, leeftijd en opleidingsniveau, opgenomen om de steekproef in kaart te brengen en de interpretatie van de bevindingen van context te kunnen voorzien. Geslacht werd gemeten met één vraag waarbij participanten konden kiezen uit de antwoordcategorieën ‘Man’, ‘Vrouw’ of ‘Anders’. Leeftijd werd gevraagd met een open invoervraag, waarbij participanten hun leeftijd in jaren opgaven in antwoord op de vraag: “Wat is je leeftijd in jaren?” Het opleidingsniveau is eveneens met één item gemeten, aan de hand van meerkeuzeantwoordopties op de vraag: “Wat is jouw hoogst voltooide opleiding?”, met antwoordcategorieën uiteenlopend van ‘Primair onderwijs (basisschool en speciaal onderwijs)’ tot ‘Doctoraat (PhD)’. De indeling van opleidingsniveau was gebaseerd op de ISCED 2011-classificatie (UNESCO, 2012) en aangepast aan de Nederlandse onderwijscontext (Ministerie van Onderwijs, Cultuur en Wetenschap, 2021)

3.5 Procedure

De dataverzameling vond plaats via een online survey. Respondenten ontvingen een bericht met daarin een link naar het onderzoek. Bij het openen van de vragenlijst kregen de respondenten eerst informatie over het doel van dit onderzoek, de geschatte duur, de vrijwilligheid van deelname en de anonieme verwerking van de gegevens. Vervolgens moesten zij met een informed consent expliciet toestemming geven voor deelname. Indien een respondent niet akkoord ging met deelname, werd deze verwezen naar de eindpagina van

het onderzoek. Na het geven van toestemming werd de respondenten gevraagd om enkele vragen te beantwoorden met betrekking tot hun demografische gegevens.

Na dit inleidende gedeelte van het experiment kregen de respondenten een korte introductietekst te lezen. In deze introductie tekst werd uitgelegd dat zij online een advertentie voor een crème van een huidverzorgingsmerk te zien zouden krijgen en dat deze advertentie met behulp van AI was gemaakt. Vervolgens kregen zij de instructie om de advertentie zorgvuldig te bekijken en een indruk te vormen van zowel de advertentie als het merk.

Na het lezen van de introductietekst kregen respondenten één advertentie te zien, afhankelijk van de conditie waaraan zij waren toegewezen. Respondenten werden door randomisatie willekeurig toegewezen aan één van de vier condities. Deze randomisatie zorgt ervoor dat verstoringende achtergrondkenmerken zoveel mogelijk gelijk worden verdeeld over de condities, waardoor systematische verschillen tussen groepen worden beperkt. Op deze manier zorgt randomisatie ervoor dat de interne validiteit van het experiment wordt vergroot (Albright & Malloy, 2000). De vier condities bestonden uit een combinatie van reclameappeal (emotioneel versus functioneel) en merkbekendheid (bekend merk versus fictief merk).

Respondenten bekeken de advertentie minimaal vijf seconden voordat zij doorgingen naar de vragenlijst. Eerst werden er vragen gesteld over de moderator van dit onderzoek, merkbekendheid. Tot slot werden er een aantal vragen gesteld over de afhankelijke variabele merkbetrouwbaarheid. Om de onafhankelijkheid van de antwoorden te waarborgen, was het voor respondenten niet mogelijk om terug te scrollen naar eerder beantwoorde vragen of de advertentie opnieuw te bekijken tijdens het invullen van de schalen. Het experiment nam naar verwachting ongeveer 5 minuten in beslag; de vragenlijst kon in één sessie worden ingevuld. Deelname aan het onderzoek was volledig vrijwillig en respondenten konden op ieder moment stoppen. Na afloop kregen zij een korte debriefing te zien waarin werd toegelicht dat het onderzoek zich richtte op de beoordeling van AI-gegenereerde reclame. Tot slot werden de respondenten bedankt voor hun deelname aan het onderzoek.

3.6 Data-analyse

Als eerste stap van de data-analyse werd de dataset van Qualtrics geëxporteerd naar IBM SPSS (versie 29.0.2) en aangevuld met de PROCESS-macro (versie 4.2). De dataset werd vervolgens opgeschoond en voorbereid op de analyse. Respondenten met een leeftijd buiten de doelgroep werden verwijderd, evenals respondenten met onvolledige antwoorden.

Vervolgens werden op basis van demografische kenmerken enkele beschrijvende analyses uitgevoerd om de steekproef te karakteriseren.

De betrouwbaarheid van de schaal voor merkbetrouwbaarheid werd bepaald door middel van Cronbachs alfa. Een waarde van .80 of hoger werd hierbij beschouwd als goed (Field, 2018). Voor merkbetrouwbaarheid werd daarbij een exploratieve factoranalyse (EFA) uitgevoerd, om de eendimensionaliteit van het construct te verifiëren. Omdat dit construct door meerdere items werd gemeten, was een EFA aangewezen om te controleren of alle items op één onderliggende factor laadden. Op basis van deze analyses zijn samengestelde schaalscores berekend door het gemiddelde van de bijbehorende items te nemen. Hogere scores stonden daarbij steeds voor een sterkere aanwezigheid van het gemeten concept.

Voorafgaand aan het toetsen van de hypothesen werd gecontroleerd of de data voldeden aan de assumpties van regressie- en variantieanalyse. Hierbij werd gekeken naar de aanwezigheid van outliers, de normaliteit van de residuen (Shapiro-Wilk-toets) de homogeniteit van varianties (Levene's toets) en is de dataset gecontroleerd op multicollineariteit middels de Variance Inflation Factor (VIF). Gezien het experimentele design met gerandomiseerde condities werden hierbij geen problematische correlaties tussen de onafhankelijke variabelen verwacht.

Om de hypothesen te toetsen werd gebruikgemaakt van de PROCESS-macro (Model 1) van Hayes (2022). Hoewel een Two-way ANOVA passend is voor een 2x2-design, biedt de regressie-gebaseerde benadering van Hayes belangrijke voordelen voor dit onderzoek. Ten eerste maakt PROCESS een meer gedetailleerde analyse van het interactie-effect mogelijk door middel van 'simple slopes'-analyses, waardoor exact kan worden bepaald onder welke specifieke condities van merkbekendheid de reclameappeal effect sorteert. PROCESS voert de analyses rechtstreeks binnen één model uit. Ten tweede biedt deze methode meer flexibiliteit om eventuele covariaten (zoals sociodemografische kenmerken) op te nemen in het model om de robuustheid van de effecten te toetsen. De onafhankelijke variabelen werden hierbij gecodeerd als dummy-variabelen. Voor alle analyses werd een significantieniveau van $p < .05$ gehanteerd. Reclameappeal en merkbekendheid zijn voorafgaand aan de moderatieanalyse gecentreerd om de interpretatie van de hoofdeffecten in het interactiemodel te verbeteren.

H4 – Resultaten

In dit hoofdstuk volgt een bespreking van de resultaten van het online experiment. Aan bod komen de beschrijvende gegevens, check op aannames, correlaties, hypothesetoetsing en tot slot de additionele analyse.

4.1 Schaalconstructie Merkbetrouwbaarheid

Voorafgaand aan de hypothesetoetsing is gecontroleerd of de tien items voor merkbetrouwbaarheid samen één schaal vormden. Eén negatief geformuleerd item is voorafgaand aan de analyse omgepooled. De data bleken geschikt voor factoranalyse, KMO = .815, en Bartlett's test of sphericity was significant, $\chi^2(45) = 279.02$, $p < .001$. Vervolgens is een exploratieve factoranalyse uitgevoerd met Principal Axis Factoring, waarbij een éénfactoroplossing werd geëxtraheerd. De eerste factor had een eigenwaarde van 4.302 en verklaarde na extractie 37.21% van de variantie. Alle tien items laadden positief op deze factor, met factorladingen tussen .430 en .770. Tegelijkertijd moet deze eendimensionaliteit voorzichtig worden geïnterpreteerd, omdat de geëxtraheerde factor 37.21% van de variantie verklaarde en daarmee een aanzienlijk deel van de variantie onverklaard bleef. Op basis van de positieve factorladingen, de goede interne consistentie en de theoretische samenhang tussen de items werd de schaal desondanks als voldoende eendimensionaal beschouwd. De interne consistentie van de schaal was goed, Cronbach's $\alpha = .839$. Daarom zijn de tien items samengevoegd tot een gemiddelde schaalscore voor merkbetrouwbaarheid, waarbij hogere scores wijzen op een hogere perceptie van merkbetrouwbaarheid.

4.2 Beschrijvende gegevens

De uiteindelijke steekproef bestond uit 143 deelnemers, waarmee de steekproef ruim voldeed aan de vooraf vastgestelde minimale steekproefomvang van 128 respondenten op basis van de poweranalyse. De 143 deelnemers werden gelijk verdeeld over de vier condities ($n = 37$ per conditie met uitzondering van de conditie met een onbekend merk en een functioneel appeal, hier was $n = 32$). De gemiddelde leeftijd van de deelnemers was 21 jaar ($M = 20.99$, $SD = 2.22$), met een minimum van 17 en een maximum van 27 jaar. Wat betreft geslacht identificeert 52.4% van de participanten zich als vrouw ($n = 75$), 46.9% als man ($n = 67$) en 0.7% als anders ($n = 1$).

Met betrekking tot het hoogst afgeronde opleidingsniveau gaf 2.8% van de participanten aan onderwijs te hebben afgerond op vmbo-niveau, praktijkonderwijs, mbo-entreeopleiding of in de onderbouw van het voortgezet onderwijs. Daarnaast gaf 62.9% aan onderwijs te hebben afgerond op havo-/vwo-niveau, vavo op havo-/vwo-niveau of mbo-

niveau 2 tot en met 4. Een klein deel van de steekproef, 1.4%, had een Associate degree afgerond. Verder gaf 27.3% aan een wo-bachelor, hbo-bachelor of post-hbo opleiding te hebben afgerond. Tot slot gaf 5.6% aan een hbo- of wo-master te hebben afgerond.

De gemiddelde merkbetrouwbaarheid na blootstelling aan een AI-gegenereerde advertentie bedroeg $M = 3.16$ ($SD = 0.62$). De gemiddelde merkbetrouwbaarheid was het hoogst in de conditie met een functionele advertentie voor een bekend merk ($M = 3.46$, $SD = 0.41$) en het laagst in de conditie met een emotionele advertentie voor een fictief merk ($M = 2.78$, $SD = 0.70$). In de conditie met een functionele advertentie voor een fictief merk bedroeg de gemiddelde merkbetrouwbaarheid $M = 3.12$ ($SD = 0.56$), terwijl de conditie met een emotionele advertentie voor een bekend merk een gemiddelde score van $M = 3.26$ ($SD = 0.57$) liet zien. Zie Tabel 3 voor de gemiddelden en standaarddeviaties van merkbetrouwbaarheid per experimentele conditie.

Tabel 3

Gemiddelden en standaarddeviaties van merkbetrouwbaarheid per conditie

	<i>M</i>	<i>SD</i>
Functionele advertentie x Bekend merk	3.46	0.41
Functionele advertentie x Fictief merk	3.12	0.56
Emotionele advertentie x Bekend merk	3.26	0.57
Emotionele advertentie x Fictief merk	2.78	0.70

Noot. Scores zijn gemeten op een vijfpuntsschaal, waarbij hogere scores wijzen op een hogere merkbetrouwbaarheid.

4.3 Check op aannames

Voorafgaand aan de hypothesetoetsing zijn de assumpties van de regressieanalyse gecontroleerd. De *lineariteit* is beoordeeld aan de hand van een scatterplot van de gestandaardiseerde residuen tegen de voorspelde waarden. Omdat de predictoren categorisch waren, lagen de punten in verticale clusters. De scatterplot liet geen duidelijke kromming of systematisch patroon zien, wat erop wijst dat de assumptie van lineariteit voldoende werd vervuld.

Multicollineariteit vormde geen probleem. De tolerance-waarden lagen tussen .325 en .481 en de VIF-waarden tussen 2.081 en 3.081. Deze waarden liggen binnen de conventionele grenswaarden van $\text{Tolerance} > .20$ en $\text{VIF} < 5$.

De *normaliteit* van de residuen werd gecontroleerd met de Shapiro-Wilk-test en de Kolmogorov-Smirnov-test. De Shapiro-Wilk-test was niet significant, $W = .989, p = .302$, wat erop wijst dat de residuen niet significant afwijken van een normale verdeling. De Kolmogorov-Smirnov-test was wel significant, $D = .082, p = .021$. Omdat deze toets bij grotere steekproeven relatief gevoelig is voor kleine afwijkingen, werd de normaliteitsassumptie op basis van de Shapiro-Wilk-test als voldoende vervuld beschouwd.

De *homogeniteit* van varianties is getoetst met Levene's test. Deze test was significant, $F(3, 139) = 3.46, p = .018$, wat betekent dat de varianties tussen de experimentele condities ongelijk waren. Daarom is bij de moderatieanalyse gebruikgemaakt van heteroscedasticity-consistent standaardfouten, specifiek de HC3-correctie.

4.4 Covariaten

Om te bepalen of geslacht als covariaat moest worden opgenomen, is een onafhankelijke t-test uitgevoerd voor het verschil in merkbetrouwbaarheid tussen mannen en vrouwen. De ene respondent die zichzelf niet als man of vrouw identificeerde, is enkel voor deze analyse buiten beschouwing gelaten, omdat het aantal te klein was voor een betrouwbare vergelijking. Mannen rapporteerden gemiddeld een merkbetrouwbaarheid van $M = 3.16$ ($SD = 0.61$), terwijl vrouwen gemiddeld eveneens een merkbetrouwbaarheid van $M = 3.16$ ($SD = 0.63$) rapporteerden. Dit verschil was niet significant, $t(140) = 0.01, p = .990, d = 0.002$. Geslacht is daarom niet als covariaat opgenomen in de hoofdanalyse.

Om vervolgens te bepalen of opleidingsniveau als covariaat moest worden opgenomen, is een one-way ANOVA uitgevoerd met merkbetrouwbaarheid als afhankelijke variabele en opleidingsniveau als factor. Levene's test was niet significant, $F(4, 138) = 0.48, p = .751$, wat erop wijst dat de assumptie van homogeniteit van varianties niet werd geschonden. De ANOVA liet geen significant verschil zien in merkbetrouwbaarheid tussen opleidingsgroepen, $F(4, 138) = 2.15, p = .078, \eta^2 = .059$. Opleidingsniveau is daarom niet als covariaat opgenomen in de hoofdanalyse.

4.5 Hypothesetoetsing

Voorafgaand aan de regressieanalyse zijn de dichotome predictoren als dummyvariabelen gecodeerd, waarbij reclameappeal is gecodeerd als 0 = functioneel en 1 = emotioneel, en merkbekendheid als 0 = fictief merk en 1 = bekend merk

Hypothese 1 stelt dat blootstelling aan een AI-gegenereerde advertentie met een functioneel appeal leidt tot een hogere perceptie van merkbetrouwbaarheid dan blootstelling aan een AI-gegenereerde advertentie met een emotioneel appeal. Deze hypothese wordt

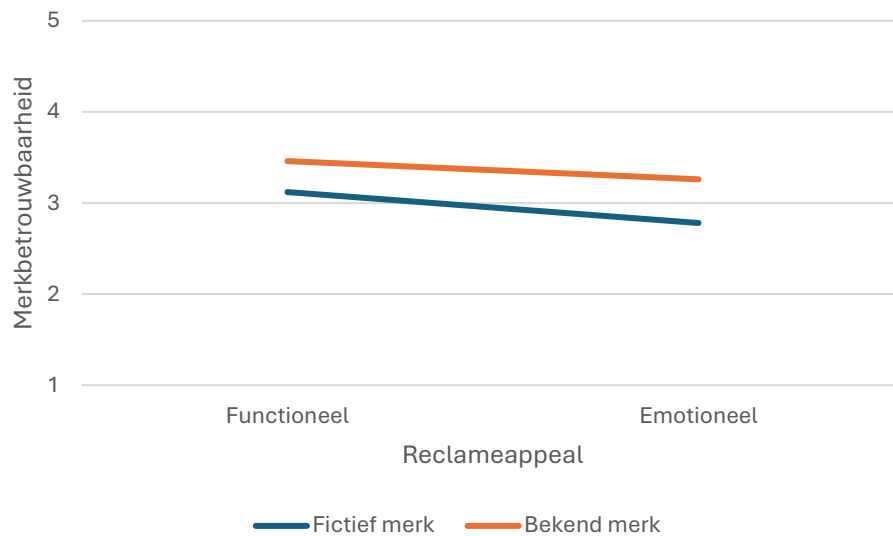
ondersteund. Reclameappeal had een significant effect op merkbetrouwbaarheid ($b = -0.265$, $SE = 0.096$, $t = -2.75$, $p = .007$, 95% BI [-0.455, -0.074]).

Omdat reclameappeal was gecodeerd als 0 = functioneel en 1 = emotioneel, betekent de negatieve coëfficiënt dat emotionele AI-gegenereerde advertenties resulteerden in een lagere merkbetrouwbaarheid dan functionele AI-gegenereerde advertenties. Dit patroon is ook zichtbaar in de beschrijvende statistieken: respondenten in de functionele reclameconditie rapporteerden gemiddeld een hogere merkbetrouwbaarheid ($M = 3.30$, $SD = 0.52$) dan respondenten in de emotionele reclameconditie ($M = 3.02$, $SD = 0.68$). Functionele AI-reclame leidde tot een hogere perceptie van merkbetrouwbaarheid dan emotionele AI-reclame.

Hypothese 2 stelt dat het positieve effect van een functioneel appeal op merkbetrouwbaarheid wordt gemodereerd door merkbekendheid, waarbij dit effect sterker is bij merken met lage merkbekendheid dan bij merken met hoge merkbekendheid. Voordat het interactie-effect werd geïnterpreteerd, is eerst gekeken naar het directe effect van merkbekendheid. Merkbekendheid had een significant positief effect op merkbetrouwbaarheid, ($b = 0.419$, $SE = 0.098$, $t = 4.28$, $p < .001$, 95% BI [0.225, 0.613]). Omdat merkbekendheid was gecodeerd als 0 = fictief merk en 1 = bekend merk, betekent dit dat het bekende merk gemiddeld als betrouwbaarder werd beoordeeld dan het fictieve merk. Dit patroon is ook zichtbaar in de beschrijvende statistieken: de gemiddelde merkbetrouwbaarheid lag hoger bij het bekende merk ($M = 3.36$, $SD = 0.51$) dan bij het fictieve merk ($M = 2.94$, $SD = 0.66$). De moderatiehypothese werd echter niet ondersteund. De interactie tussen reclameappeal en merkbekendheid was niet significant ($b = 0.135$, $SE = 0.194$, $t = 0.69$, $p = .490$, 95% BI [-0.250, 0.519]). Dit betekent dat merkbekendheid het effect van reclameappeal op merkbetrouwbaarheid niet significant modereerde. Het volledige model verklaarde 16.8% van de variantie in merkbetrouwbaarheid ($R^2 = .168$, $F(3, 139) = 9.17$, $p < .001$).

Figuur 3

Gemiddelde merkbetrouwbaarheid per reclameappeal en merkbekendheid



Noot. Scores zijn gemeten op een vijfpuntsschaal, waarbij hogere scores wijzen op een hogere perceptie van merkbetrouwbaarheid.

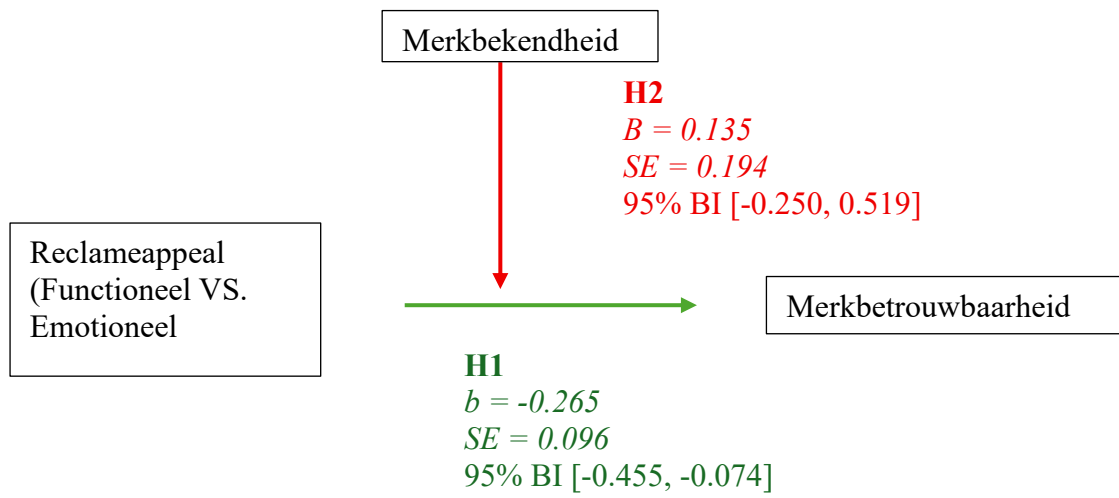
4.6 Aanvullende analyses

Als aanvulling op de hypothesetoetsing is een additionele regressieanalyse uitgevoerd waarin huidverzorgingsproducten-betrokkenheid en AI-houding als losse voorspellers van merkbetrouwbaarheid zijn opgenomen. Het model was significant, $F(2, 140) = 4.57$, $p = .012$, en verklaarde 6.1% van de variantie in merkbetrouwbaarheid, $R^2 = .061$.

Huidverzorgingsproduct-betrokkenheid bleek een significante positieve voorspeller van merkbetrouwbaarheid, $b = 0.120$, $SE = 0.052$, $\beta = .188$, $t = 2.30$, $p = .023$. Dit betekent dat participanten met een hogere betrokkenheid bij huidverzorging gemiddeld een hogere merkbetrouwbaarheid rapporteerden. De houding tegenover AI was geen significante voorspeller van merkbetrouwbaarheid, $b = 0.113$, $SE = 0.060$, $\beta = .155$, $t = 1.89$, $p = .061$, hoewel de richting van het verband positief was.

Figuur 4

Getoetst model met coëfficiënten



H5 - Conclusie en discussie

In dit hoofdstuk worden de conclusie en discussie van het experiment besproken. De discussie omvat de theoretische terugkoppeling van de bevindingen, alternatieve theoretische verklaringen voor de onderzoeksresultaten, methodologische beperkingen, overige suggesties voor vervolgonderzoek en tot slot de wetenschappelijke en praktische implicaties.

5.1 Conclusie

Dit onderzoek beantwoordt de vraag in hoeverre reclameappeal invloed heeft op de perceptie van merkbetrouwbaarheid bij AI-gegenereerde reclame onder Generatie Z en of dit effect afhankelijk is van de mate van merkbekendheid. De resultaten laten zien dat reclameappeal een significante rol speelt in de evaluatie van merkbetrouwbaarheid. Blootstelling aan een AI-gegenereerde reclame met een functioneel appeal leidde tot een hogere beoordeling van merkbetrouwbaarheid dan blootstelling aan een AI-gegenereerde reclame met een emotioneel appeal. De resultaten suggereren dat Generatie Z een merk betrouwbaarder beoordeelt als de AI-gegenereerde reclame een boodschap heeft die aansluit bij rationale, concrete en productgerichte informatie.

De verwachte modererende rol van merkbekendheid werd niet statistisch ondersteund. De verwachting was dat merkbekendheid de lagere merkbetrouwbaarheid van een emotionele AI-appeal zou kunnen bufferen, waardoor het voordeel van een functioneel appeal sterker zou zijn bij een fictief merk dan bij een bekend merk. Deze interactie werd echter niet gevonden. Het verschil in merkbetrouwbaarheid tussen functionele en emotionele AI-reclame verschilde niet significant tussen het bekende en het fictieve merk.

Aanvullend liet de analyse zien dat merkbekendheid als onafhankelijke voorspeller samenhang met merkbetrouwbaarheid. Het bekende merk werd gemiddeld als betrouwbaarder beoordeeld dan het fictieve merk. Dit suggereert dat merkbekendheid in dit onderzoek bijdroeg aan een hoger algemeen niveau van merkbetrouwbaarheid, maar niet bepaalde hoe sterk het effect van reclameappeal was.

Tot slot liet de additionele analyse zien dat betrokkenheid bij het product huidverzorging een significante positieve voorspeller was van merkbetrouwbaarheid. Participanten met een hogere betrokkenheid bij huidverzorgingsproducten beoordeelden het merk gemiddeld als betrouwbaarder. Dit betekent dat de intrinsieke interesse van de consument in de productcategorie samenhangt met een hogere algemene beoordeling van merkbetrouwbaarheid. Consumenten die veel affiniteit hebben met huidverzorging staan

mogelijk meer open voor de informatie in de advertentie, wat hun algemene betrouwbaarheidsperceptie van de afzender verhoogt, los van de technologische totstandkoming van de reclame. Het suggereert dat consumenten die de categorie 'belangrijk' vinden, minder kritisch kijken naar de technologische totstandkoming van de advertentie en meer naar de productinhoud zelf.

5.2 Theoretische terugkoppeling van de bevindingen

Dit onderzoek draagt bij aan de literatuur over AI-gegenereerde reclame door niet de vergelijking tussen mens en AI centraal te stellen, maar juist verschillen binnen AI-gegenereerde reclame te onderzoeken. Waar eerdere studies zich veelal richten op de verschillen tussen AI en mens als bron, laat dit onderzoek zien dat vooral de manier waarop AI strategisch wordt ingezet relevant is voor de beoordeling van merkbetrouwbaarheid. De resultaten laten zien dat een functioneel appeal binnen AI-gegenereerde reclame leidde tot een hogere perceptie van merkbetrouwbaarheid dan een emotioneel appeal. Dit patroon is zichtbaar in de beschrijvende statistieken: bij het fictieve merk lag de merkbetrouwbaarheid hoger in de functionele conditie ($M = 3.12$, $SD = 0.56$) dan in de emotionele conditie ($M = 2.78$, $SD = 0.70$). Ook bij het bekende merk scoorde de functionele conditie hoger ($M = 3.46$, $SD = 0.41$) dan de emotionele conditie ($M = 3.26$, $SD = 0.57$). De bevindingen suggereren daarmee dat functionele AI-reclame gunstiger wordt beoordeeld dan emotionele AI-reclame, ongeacht of de afzender een fictief of bekend merk is. Tegelijkertijd moeten deze resultaten genuanceerd worden geïnterpreteerd. Hoewel functionele appeals significant hoger scoorden dan emotionele appeals, bleven de absolute gemiddelden voor merkbetrouwbaarheid gematigd. De gemiddelde scores bewogen tussen $M = 2.78$ en $M = 3.46$ op een vijfpuntsschaal. Dit betekent dat een functionele insteek de merkbetrouwbaarheid wel lijkt te verhogen ten opzichte van een emotionele insteek, maar niet automatisch leidt tot substantieel hoge niveaus van absoluut vertrouwen.

Deze bevindingen kunnen worden geïnterpreteerd vanuit de machine heuristic, waarin algoritmen vaak worden geassocieerd met rationaliteit, objectiviteit, efficiëntie en dataverwerking (Sundar & Kim, 2019). Als een AI-gegenereerde advertentie een functionele boodschap bevat, kan er een inhoudelijke aansluiting tussen de bron van de boodschap en de boodschap zelf ontstaan. Het functionele appeal benadrukt concrete producteigenschappen en rationele voordelen, waardoor deze beter past bij het beeld van AI als een logische en datagedreven bron. Dit biedt een mogelijke verklaring voor waarom respondenten in dit

onderzoek functionele AI-gegenereerde advertenties als betrouwbaarder beoordeelden dan emotionele AI-gegenereerde advertenties.

Daarnaast kunnen de resultaten theoretisch worden geïnterpreteerd vanuit de match-up hypothesis. Deze theorie stelt dat een boodschap als overtuigender wordt ervaren wanneer er congruentie bestaat tussen de eigenschappen van de bron en de inhoud van de boodschap (Lee et al., 2022). In de context van dit onderzoek betekent dit dat een functioneel reclameappeal beter aansluit bij de gepercipieerde aard van AI dan een emotioneel reclameappeal. Een emotioneel appeal vraagt om oprechtheid, menselijke betrokkenheid en affectieve geloofwaardigheid, terwijl AI mogelijk juist wordt gezien als een algoritmische bron zonder eigen gevoelens of intrinsieke intenties. Daardoor kan een emotionele AI-gegenereerde advertentie als minder congruent worden ervaren, en heeft het dus een minder positief effect op merkbetrouwbaarheid dan een functioneel appeal. Deze interpretatie moet echter voorzichtig worden benaderd. De onderliggende psychologische mechanismen, zoals de aanwezigheid van de machine heuristic of de bewuste perceptie van congruentie tussen bron en boodschap, zijn in dit onderzoek niet rechtstreeks gemeten. Er is bijvoorbeeld geen mediatieanalyse uitgevoerd met waargenomen authenticiteit, broncredibiliteit of perceived fit als verklarende variabelen. De verklaring dat de hogere merkbetrouwbaarheid voortkomt uit gepercipieerde congruentie blijft daarom een plausibele, theoretisch gefundeerde interpretatie, maar kan op basis van dit onderzoeksdesign niet causaal worden vastgesteld. Vervolgonderzoek zou deze mechanismen expliciet moeten meten om te bepalen of perceived fit daadwerkelijk verklaart waarom functionele AI-reclame betrouwbaarder wordt beoordeeld dan emotionele AI-reclame.

Aanvullend kunnen de resultaten post-hoc worden geïnterpreteerd vanuit de source credibility theory. Deze theorie is niet vooraf als verklarend mechanisme in het theoretisch kader getoetst, maar kan wel helpen om de bevindingen nader te duiden. Deze theorie benadrukt dat de overtuigingskracht van een boodschap afhankelijk is van de waargenomen geloofwaardigheid van de bron (Hovland & Weiss, 1951; Ohanian, 1990). Hoewel AI binnen reclame mogelijk relatief sterk scoort op expertise, omdat algoritmen worden geassocieerd met data, efficiëntie en technische nauwkeurigheid, kan AI tegelijkertijd zwakker scoren op de vertrouwensdimensie. Consumenten zien algoritmes doorgaans niet als bronnen met oprechte of menselijke intenties, omdat algoritmische beslissingen vaak worden ervaren als minder intuïtief, minder sociaal en minder menselijk dan beslissingen van menselijke actoren (Lee, 2018). De resultaten van dit onderzoek suggereren dat een functioneel appeal beter

gebruik maakt van de expertise-associaties van AI, terwijl een emotioneel appeal juist de zwakkere vertrouwensdimensie van AI blootlegt.

Tot slot toont dit onderzoek aan dat merkbekendheid geen significante moderator vormt, maar wel samenhangt met een hogere perceptie van merkbetrouwbaarheid. Dit laat zien dat respondenten een bekend merk, ongeacht het gebruikte reclameappeal, over het algemeen als betrouwbaarder beoordelen dan een fictief merk. Deze bevinding sluit aan bij de literatuur over brand awareness, waarin wordt gesteld dat bekende merken sterker aanwezig zijn in het geheugen van consumenten en daardoor gemakkelijker positieve merkassociaties oproepen (Keller, 1993). Ook past dit bij eerder onderzoek waaruit blijkt dat merkbekendheid kan bijdragen aan merkvertrouwen, doordat consumenten bij bekende merken beschikken over meer eerdere ervaringen, kennis en beoordelingszekerheid (Ha & Perks, 2005; Laroche et al., 1996). Deze bevinding suggereert dat merkbekendheid in dit onderzoek eerder fungeerde als algemene vertrouwensbasis dan als moderator die de relatie tussen reclameappeal en merkbetrouwbaarheid veranderde. Daarmee wordt H2 niet ondersteund, terwijl de directe rol van merkbekendheid wel aansluit bij bestaande literatuur over brand awareness en merkvertrouwen.

Deze directe rol van merkbekendheid als onafhankelijke voorspeller kan helpen verklaren waarom het verwachte modererende effect van merkbekendheid op de relatie tussen reclameappeal en merkbetrouwbaarheid uitbleef. Mogelijk vormde merkbekendheid in dit onderzoek eerder een algemene vertrouwensbasis dan een factor die bepaalde hoe sterk reclameappeal doorwerkte in merkbetrouwbaarheid. Bekende merken werden in beide reclameappeals relatief positief beoordeeld, waardoor de invloed van het functionele of emotionele appeal minder afhankelijk werd van de mate van merkbekendheid. Tegelijkertijd bleef het effect van reclameappeal wel zichtbaar: functionele AI-gegenereerde advertenties leidden gemiddeld tot hogere merkbetrouwbaarheid dan emotionele AI-gegenereerde advertenties. Deze bevinding suggereert dat de congruentie tussen AI als rationele, datagedreven bron en een functionele boodschap een centrale rol speelt in de beoordeling van merkbetrouwbaarheid, in lijn met de machine heuristic en de match-up hypothesis (Sundar & Kim, 2019; Lee et al., 2022).

5.3 Alternatieve theoretische verklaringen voor de onderzoeksresultaten

Een eerste verklaring voor het uitblijven van het verwachte modererende effect zou kunnen zijn dat merkbekendheid in dit onderzoek vooral fungeerde als zelfstandig vertrouwens ‘anker’ en niet als een factor die invloed had op het effect van reclameappeal.

Het bekende merk werd gemiddeld als betrouwbaarder beoordeeld dan het fictieve merk, ongeacht of de advertentie een functioneel of emotioneel appeal bevatte. Dit sluit aan bij de gedachte dat bekende merken bestaande kennisstructuren, associaties en ervaringen activeren bij consumenten (Keller, 1993). Wanneer consumenten al bekend zijn met een merk, hoeven zij hun oordeel over dat merk niet enkel te baseren op één specifieke advertentie die zij te zien krijgen. Het merk zelf biedt dan al een basis voor vertrouwen. Hierdoor kan merkbekendheid wel bijdragen aan een hogere algemene merkbetrouwbaarheid, maar hoeft het niet noodzakelijk te bepalen of een functionele of emotionele AI-gegenereerde advertentie overtuigender is.

Een tweede mogelijke verklaring is dat merkbekendheid vooral fungeerde als een direct en onafhankelijke anker voor vertrouwen. Het bekende merk werd gemiddeld betrouwbaarder beoordeeld dan het fictieve merk, ongeacht het gebruikte reclameappeal. Dit suggereert dat de bestaande merkassociaties rondom Nivea wellicht een algemene positieve basis vormden voor de beoordeling van merkbetrouwbaarheid, zonder dat deze associaties de verwerking van functionele versus emotionele AI-reclame significant veranderden. De gemiddelden wijzen in deze richting: bij het fictieve merk was het verschil tussen de functionele advertentie ($M = 3.12$, $SD = 0.56$) en de emotionele advertentie ($M = 2.78$, $SD = 0.70$) groter dan bij het bekende merk, waar de functionele advertentie ($M = 3.46$, $SD = 0.41$) en de emotionele advertentie ($M = 3.26$, $SD = 0.57$) dicht bij elkaar lagen. Dit descriptieve patroon ligt in de voorspelde richting van de verwachte moderatie, maar de interactieterm was niet significant. Daarom biedt dit patroon geen statistisch bewijs voor de verwachte bufferende werking van merkbekendheid.

5.4 Methodologische beperkingen

Een eerste beperking heeft betrekking op de operationalisering van merkbetrouwbaarheid. Merkbetrouwbaarheid is gemeten met items gebaseerd op de Brand Trust Scale van Munuera-Aleman et al. (2003). Deze schaal is geschikt om vertrouwen in merken te meten, maar een deel van de items veronderstelt een concrete klacht/gebruik situatie. Het zou mogelijk kunnen zijn dat een deel van de items inhoudelijk minder goed aansloot op de experimentele context van dit onderzoek. Daardoor kan het voor respondenten lastig zijn geweest om stellingen te beantwoorden over bijvoorbeeld de mate waarin een merk problemen zou oplossen, consumenten tevreden zou stellen of compensatie zou bieden. Dit kan ertoe hebben geleid dat respondenten hun antwoorden baseerden op algemene indrukken van betrouwbaarheid of sympathie, in plaats van op daadwerkelijk vertrouwen in het merk.

Hierdoor moet voorzichtig worden omgegaan met de interpretatie van merkbetrouwbaarheid als betrouwbare ontwikkelde brand trust-construct.

Daarnaast is merkbetrouwbaarheid in dit onderzoek gemeten na één enkele advertentieblootstelling. Merkbetrouwbaarheid is een construct dat zich ontwikkelt op basis van herhaalde ervaringen met een merk, eerdere interacties en bijvoorbeeld productgebruik. Ha en Perks (2005) laten zien dat brand trust wordt gevormd door merkervaringen, merkbekendheid en klanttevredenheid, terwijl Motta-Filho (2017) benadrukt dat consistente klantervaringen het vertrouwen in het vermogen van een merk versterken om verwachtingen waar te maken. In dit onderzoek werd merkbetrouwbaarheid gemeten direct nadat de respondenten een van de AI-gegenereerde advertenties kort hadden bekeken. Hierdoor is het mogelijk dat de meting vooral de eerste indruk van merkbetrouwbaarheid weerspiegelt in plaats van diepgeworteld vertrouwen in het merk. Dit is met name relevant voor het fictieve merk Luxera, waarbij respondenten nog geen bestaande ervaringen of associaties hadden waarop zij hun oordeel konden baseren. De resultaten moeten daarom worden geïnterpreteerd als effecten op waargenomen merkbetrouwbaarheid in een eerste advertentiecontact, en niet als effecten op langdurig merkvertrouwen.

Een derde beperking van dit onderzoek gaat over de operationalisering van merkbekendheid. In dit onderzoek werd merkbekendheid gemanipuleerd door een bekend merk, Nivea, te vergelijken met een fictief merk, Luxera. Deze keuze zorgde voor een duidelijk contrast tussen hoge en lage merkbekendheid. In de praktijk is merkbekendheid echter geen dichotoom construct. Door merkbekendheid te operationaliseren als bekend versus fictief merk blijft onduidelijk welk aspect van merkbekendheid precies verantwoordelijk was voor de gevonden verschillen. Bovendien verschilt Nivea mogelijk niet alleen van Luxera in bekendheid, maar ook in specifieke merkassociaties zoals zorg, huidvriendelijkheid, kwaliteit of nostalgie. Hierdoor kan niet volledig worden uitgesloten dat het effect van merkbekendheid deels werd beïnvloed door merk specifieke associaties van Nivea.

Tot slot de gekozen productcategorie: in dit onderzoek is gekozen voor een huidverzorgingsproduct, crème. Deze productcategorie is zeer relevant voor het onderzoeken van merkbetrouwbaarheid, omdat consumenten binnen deze productcategorie merkbetrouwbaarheid erg belangrijk vinden (Sanny et al., 2020). Tegelijkertijd kan deze productcontext het effect van reclameappeal hebben beïnvloed. Bij een product dat direct op het lichaam wordt gebruikt, kunnen consumenten sterker letten op concrete informatie over werking, ingrediënten of betrouwbaarheid. Hierdoor kan een functioneel appeal relatief

overtuigend zijn geweest, terwijl een emotioneel appeal minder goed aansloot bij de informatiebehoefte rondom dit product. De bevinding dat functionele AI-gegenereerde reclame tot hogere merkbetrouwbaarheid leidde dan emotionele AI-gegenereerde reclame moet daarom voorzichtig worden gegeneraliseerd naar andere productcategorieën. Bij producten waarbij beleving, identiteit of symbolische waarde centraler staan, zoals parfum, mode of entertainment, kan een emotionele AI-gegenereerde advertentie mogelijk anders worden beoordeeld.

5.5 Suggesties voor vervolgonderzoek

Op basis van de bevindingen en de beperkingen van dit onderzoek zijn er verschillende mogelijkheden voor vervolgonderzoek. Een eerste suggestie is om merkbetrouwbaarheid over meerdere contactmomenten te onderzoeken. In dit onderzoek werd merkbetrouwbaarheid gemeten na één blootstelling aan een AI-gegenereerde advertentie. Voor vervolgonderzoek is het interessant om te onderzoeken wat er met het effect gebeurt na meerdere contactmomenten. Wordt het effect van reclameappeal sterker, zwakker of stabiel wanneer consumenten meerdere AI-gegenereerde advertenties van hetzelfde merk te zien krijgen? Dit is relevant, omdat merkvertrouwen doorgaans niet ontstaat op basis van één contactmoment, maar zich ontwikkelt door herhaalde ervaringen, consistente communicatie en eerdere interacties met een merk. Een longitudinaal onderzoek kan daardoor meer inzicht bieden in de vraag of functionele AI-reclame vooral een eerste betrouwbaarheidsindruk versterkt, of ook op langere termijn bijdraagt aan merkvertrouwen.

Een tweede interessante richting voor vervolgonderzoek ligt in de operationalisering van merkbekendheid. In dit onderzoek werd merkbekendheid gemanipuleerd door een bekend merk te vergelijken met een fictief merk. Deze opzet creëerde een duidelijk contrast, maar maakte het moeilijker om onderscheid te maken tussen verschillende niveaus van merkbekendheid. Vervolgonderzoek zou merkbekendheid kunnen meten als continu construct. Denk hierbij bijvoorbeeld aan onderscheid maken tussen merkherkenning, merkherinnering, eerdere gebruikservaring en affectieve merkassociaties. Op deze manier kan worden onderzocht of er een kantelpunt bestaat waarop merkbekendheid de beoordeling van AI-gegenereerde reclame sterker begint te beïnvloeden. Ook kan vervolgonderzoek meerdere bekende en onbekende merken opnemen in het onderzoeksdesign. Zo zou merkbekendheid beter kunnen worden losgekoppeld van merk specifieke associaties.

Een derde suggestie is om het onderzoek te herhalen binnen andere productcategorieën. De huidige studie richtte zich op huidverzorging, een productcategorie

waarin veiligheid, werking en betrouwbaarheid belangrijk zijn. Daardoor kan een functioneel appeal relatief overtuigend zijn geweest. Bij productcategorieën waarin beleving, identiteit of symbolische waarde sterker centraal staan, zoals parfum, mode of entertainment, kan een emotionele AI-gegenereerde advertentie mogelijk anders worden beoordeeld.

Vervolgonderzoek kan daarom producttype als extra factor opnemen en vergelijken of het voordeel van functionele AI-reclame vooral optreedt bij functionele producten, of ook zichtbaar is bij meer expressieve en emotionele productcategorieën.

5.6 Wetenschappelijke implicaties

Dit onderzoek levert een bijdrage aan de academische literatuur over AI-gegenereerde reclame. Waar bestaand onderzoek zich vaak richt op de vergelijking tussen menselijke makers en AI, verschuift dit onderzoek de focus naar verschillen binnen het strategische gebruik van AI-gegenereerde reclame. Daarmee laat deze studie zien dat het niet alleen relevant is om verschillen tussen mens en AI te onderzoeken, maar ook om te onderzoeken welke vormen van AI-gegenereerde reclame meer of minder merkbetrouwbaarheid oproepen onder Generatie Z.

De bevinding dat functionele AI-gegenereerde advertenties tot hogere merkbetrouwbaarheid leiden dan emotionele AI-gegenereerde advertenties biedt een empirische aanvulling op theorieën zoals de machine heuristic (Sundar & Kim, 2019) en de match-up hypothesis (Lee et al., 2022). De resultaten zijn in lijn met de gedachte dat een functionele boodschap beter kan aansluiten bij de rationele en datagedreven eigenschappen die consumenten mogelijk aan AI toeschrijven. Tegelijkertijd moet deze interpretatie voorzichtig worden benaderd, omdat de onderliggende percepties van AI, zoals waargenomen objectiviteit, authenticiteit of perceived fit, in dit onderzoek niet rechtstreeks zijn gemeten.

Daarnaast draagt dit onderzoek bij aan theorievorming over merkbekendheid in een gedigitaliseerd reclamelandschap. Merkbekendheid bleek geen significante moderator van het effect van reclameappeal op merkbetrouwbaarheid, maar hing wel direct samen met een hogere perceptie van merkbetrouwbaarheid. Dit suggereert dat een gevestigde merknaam vooral kan fungeren als een algemeen betrouwbaarheidsanker, zonder dat dit noodzakelijk betekent dat merkbekendheid de verwerking van functionele versus emotionele AI-reclame verandert. Daarmee biedt dit onderzoek een nuancering van de bufferende werking van brand equity: merkbekendheid kan bijdragen aan hogere algemene merkbetrouwbaarheid, maar is in deze studie niet voldoende om de relatieve werking van verschillende AI-appeals significant te veranderen.

5.7 Praktische implicaties

De resultaten van dit onderzoek bieden richtlijnen voor reclamemakers en merken die generatieve AI willen inzetten in commerciële communicatie.

In dit onderzoek leidde een functionele AI-gegenereerde advertentie tot een hogere perceptie van merkbetrouwbaarheid dan een emotionele AI-gegenereerde advertentie. Voor marketeers suggereert dit dat AI vooral effectief kan zijn wanneer de technologie wordt ingezet voor reclameboodschappen waarin concrete productinformatie, rationele voordelen en functionele eigenschappen centraal staan. Omdat H2 niet werd ondersteund lijkt dit advies te gelden voor zowel onbekende merken als bekende merken. Voor adverteerders is het verstandig om AI-gegenereerde reclame in eerste instantie functioneel en productgericht vorm te geven, ongeacht de mate van merkbekendheid.

Voor merken betekent dit dat zij voorzichtig moeten zijn met het inzetten van AI voor sterk emotionele merkcommunicatie. Emotionele appeals vragen om oprechtheid, menselijke betrokkenheid en authenticiteit. Wanneer consumenten weten dat een emotionele advertentie met behulp van AI is ontwikkeld, kan dit de geloofwaardigheid van de boodschap verlagen. Dit betekent niet dat AI ongeschikt is voor alle vormen van creatieve of emotionele reclame, maar wel dat merken kritisch moeten nadenken over de congruentie tussen de AI-bron en de inhoud van de boodschap. Wanneer AI wordt gebruikt voor emotionele reclame, kan het belangrijk zijn om menselijke controle, creatieve supervisie of de rol van menselijke makers expliciet te maken. Vooral bij nieuwe en onbekende merken lijkt het belangrijk om terughoudend te zijn met emotionele AI-advertenties, dit omdat er nog geen sterke vertrouwensbasis of bestaande merkassociaties aanwezig zijn waarop consumenten hun beoordeling kunnen baseren.

Daarnaast laten de resultaten zien dat het bekende merk gemiddeld als betrouwbaarder werd beoordeeld dan het fictieve merk, maar dat merkbekendheid het effect van reclameappeal niet significant veranderde. Voor de praktijk betekent dit dat de voordelen van een functionele AI-appeal niet uitsluitend lijken voorbehouden aan gevestigde merken. Ook minder bekende merken kunnen baat hebben bij AI-gegenereerde advertenties waarin de nadruk ligt op duidelijke, concrete en productgerichte informatie. Tegelijkertijd blijft merkbekendheid wel relevant, omdat bekende merken in het algemeen betrouwbaarder werden beoordeeld. Voor minder bekende merken is het daarom belangrijk om niet alleen een passend reclameappeal te kiezen, maar ook actief te investeren in het opbouwen van herkenbaarheid, consistentie en vertrouwen.

Tot slot zijn de resultaten relevant voor de manier waarop merken Generatie Z willen bereiken. Deze doelgroep is opgegroeid met digitale technologie, maar blijft kritisch ten aanzien van authenticiteit en commerciële intenties. Voor deze doelgroep lijkt het daarom belangrijk dat AI-gebruik niet wordt ingezet om menselijke emotie kunstmatig na te bootsen, maar juist om duidelijke, relevante en geloofwaardige informatie te versterken.

Reclamemakers die AI gebruiken, doen er daarom goed aan om niet enkel te focussen op het feit of een advertentie technisch goed of visueel aantrekkelijk is, maar vooral of de gekozen boodschap past bij de manier waarop consumenten AI als bron interpreteren. De bevindingen laten zien dat functionele AI-reclame betrouwbaarder wordt beoordeeld dan emotionele AI-reclame, terwijl merkbekendheid deze relatie niet significant verandert. Daarmee nuanceert dit onderzoek de aanname dat Generatie Z AI-reclame vanzelfsprekend accepteert. In een communicatielandschap waarin generatieve AI steeds vanzelfsprekender wordt, lijkt vertrouwen niet primair te ontstaan door technologische vernieuwing zelf, maar door de mate waarin merken technologie strategisch, geloofwaardig en inhoudelijk passend inzetten.

Literatuurlijst

- Albright, L., & Malloy, T. E. (2000). Experimental Validity: Brunswik, Campbell, Cronbach, and Enduring Issues. *Review Of General Psychology*, 4(4), 337–353. <https://doi.org/10.1037/1089-2680.4.4.337>
- Bachnik, K., Nowacki, R., Fandrejewska-Nowakowska, A., Zatwarnicka-Madura, B., & Eisenbardt, T. (2025). AI-Generated Advertising Through the Eyes of Generation Z. *Contemporary Economics*, 417. <https://doi.org/10.5709/ce.1897-9254.575>
- Faul, F., Erdfelder, E., Buchner, A., & Lang, A.-G. (2009). Statistical power analyses using G*Power 3.1: Tests for correlation and regression analyses. *Behavior Research Methods*, 41, 1149-1160.
- Field, A. (2018). *Discovering statistics using IBM SPSS statistics* (5e ed.). Sage.
- Ha, H., & Perks, H. (2005). Effects of consumer perceptions of brand experience on the web: brand familiarity, satisfaction and brand trust. *Journal Of Consumer Behaviour*, 4(6), 438–452. <https://doi.org/10.1002/cb.29>
- Hayes, A. F. (2022). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (3rd ed.). The Guilford Press.
- Harvard Business Impact Education. (2026). <https://hbsp.harvard.edu/inspiring-minds/how-gen-z-is-using-ai>
- Have, R. T. (2026, 8 maart). *AI Marketing Statistieken 2026 | 50+ Cijfers & Data*. Searchlab. <https://searchlab.nl/statistieken/ai-marketing-statistieken-2026>
- Henke, J. (2025). The new normal: The increasing adoption of generative AI in university communication. *Journal Of Science Communication*, 24(2). <https://doi.org/10.22323/2.24020207>

- Hovland, C. I., & Weiss, W. (1951). The Influence of Source Credibility on Communication Effectiveness. *Public Opinion Quarterly*, 15(4), 635. <https://doi.org/10.1086/26635>
- IAB. (2024, 9 december). *IAB | Why young consumers aren't yet buying into Gen AI ads*. <https://www.iab.com/insights/why-young-consumers-avoid-gen-ai-ads/#:~:text=toward%20a%20more%20positive%20viewpoint,as%20the%20technology%20matures>
- Keller, K. L. (1993). Conceptualizing, Measuring, and Managing Customer-Based Brand Equity. *Journal Of Marketing*, 57(1), 1. <https://doi.org/10.2307/1252054>
- Laroche, M., Kim, C., & Zhou, L. (1996). Brand familiarity and confidence as determinants of purchase intention: An empirical test in a multiple brand context. *Journal Of Business Research*, 37(2), 115–120. [https://doi.org/10.1016/0148-2963\(96\)00056-2](https://doi.org/10.1016/0148-2963(96)00056-2)
- Lee, J., Chang, H., & Zhang, L. (2022). An integrated model of congruence and credibility in celebrity endorsement. *International Journal Of Advertising*, 41(7), 1358–1381. <https://doi.org/10.1080/02650487.2021.2020563>
- Lee, M. K. (2018). Understanding perception of algorithmic decisions: Fairness, trust, and emotion in response to algorithmic management. *Big Data & Society*, 5(1). <https://doi.org/10.1177/2053951718756684>
- Leonidou, L. C., & Leonidou, C. N. (2009). Rational Versus Emotional Appeals in Newspaper Advertising: Copy, Art, and Layout Differences. *Journal Of Promotion Management*, 15(4), 522–546. <https://doi.org/10.1080/10496490903281353>
- Morales, A. C., Amir, O., & Lee, L. (2017). Keeping It Real in Experimental Research—Understanding When, Where, and How to Enhance Realism and Measure Consumer Behavior. *Journal Of Consumer Research*, 44(2), 465–476. <https://doi.org/10.1093/jcr/ucx048>

Motta-Filho, M. (2017). Building brand trust through customers' experiences. In *Edward Elgar Publishing eBooks*. <https://doi.org/10.4337/9781785369483.00022>

Ministerie van Onderwijs, Cultuur en Wetenschap. (2021). Toedeling Nederlandse onderwijsprogramma's (ISCED). Ocwincijfers. Geraadpleegd op 10 maart 2025, <https://ap.lc/zglGo>

Munuera-Aleman, J. L., Delgado-Ballester, E., & Yague-Guillen, M. J. (2003). Development and Validation of a Brand Trust Scale. *International Journal Of Market Research*, 45(1), 1–18. <https://doi.org/10.1177/147078530304500103>

Ohanian, R. (1990). Construction and Validation of a Scale to Measure Celebrity Endorsers' Perceived Expertise, Trustworthiness, and Attractiveness. *Journal Of Advertising*, 19(3), 39–52. <https://doi.org/10.1080/00913367.1990.10673191>

Orne, M. T. (1962). On the social psychology of the psychological experiment: With particular reference to demand characteristics and their implications. *American Psychologist*, 17(11), 776–783. <https://doi.org/10.1037/h0043424>

Persada, S. F., Miraja, B. A., & Nadlifatin, R. (2019). Understanding the Generation Z Behavior on D-Learning: A Unified Theory of Acceptance and Use of Technology (UTAUT) Approach. *International Journal Of Emerging Technologies in Learning (iJET)*, 14(05), 20. <https://doi.org/10.3991/ijet.v14i05.9993>

Sanny, L., Arina, A. N., Maulidya, R. T., & Pertiwi, R. P. (2020). Purchase intention on Indonesia male's skin care by social media marketing effect towards brand image and brand trust. *Management Science Letters*, 2139–2146. <https://doi.org/10.5267/j.msl.2020.3.023>

Song, M., Chen, H., Wang, Y., & Duan, Y. (2024). Can AI fully replace human designers? Matching effects between declared creator types and advertising appeals on tourists'

- visit intentions. *Journal Of Destination Marketing & Management*, 32, 100892. <https://doi.org/10.1016/j.jdmm.2024.100892>
- Segura, E. (2024, 5 september). Gen Z: Meesters in technologie, maar tegen welke prijs? Mastercard. <https://www.mastercard.com/nl/nl/news-and-trends/perspectives/2024/gen-z-and-technology-the-impact-on-workplace-readiness.html>
- Sundar, S. S., & Kim, J. (2019). Machine heuristic: When we trust computers more than humans with our personal information. *https://dl.acm.org/doi/proceedings/10.1145/3290605*, 1–9. <https://doi.org/10.1145/3290605.3300768>
- UNESCO Institute for Statistics. (2012). International Standard Classification of Education (ISCED) 2011. <https://doi.org/10.15220/978-92-9189-123-8-en>
- Weatherbed, J. (2026, 24 januari). Get ready for the AI ad-pocalypse. *The Verge*. <https://www.theverge.com/report/866775/ai-generated-ads-slop-human-creativity>
- Wu, L., & Wen, T. J. (2021). Understanding AI Advertising From the Consumer Perspective. *Journal Of Advertising Research*, 61(2), 133–146. <https://doi.org/10.2501/jar-2021-004>
- Zajonc, R. B. (1968). Attitudinal effects of mere exposure. *Journal Of Personality And Social Psychology*, 9(2, Pt.2), 1–27. <https://doi.org/10.1037/h0025848>

Bijlage A – Informatiebrief en toestemmingsformulier

Informatie over het onderzoek

Bedankt voor uw interesse in deze studie. Het doel van dit onderzoek is om meer inzicht te krijgen in hoe consumenten advertenties ervaren en beoordelen. Hier gaan we u een aantal vragen over stellen. In deze enquête zullen we u vragen een aantal vragen te beantwoorden. Het onderzoek duurt ongeveer 5 minuten. Opslaan en later doorgaan is niet mogelijk, vul daarom de enquête in één keer in. Daarnaast willen we u vragen om de enquête in een rustige omgeving in te vullen met minimale afleiding. Tot slot benadrukken we hier dat u alleen aan de enquête mee kunt doen als u **16 jaar of ouder bent**.

Informatie over het gebruik van de data

Voor het uitvoeren van het onderzoek is het noodzakelijk dat uw persoonsgegevens worden verzameld, gebruikt en opgeslagen. Het gaat om persoonsgegevens zoals demografische gegevens. In het toestemmingsformulier wordt u gevraagd toestemming te geven voor het verzamelen, gebruiken en opslaan van uw persoonsgegevens. Ook wordt u gevraagd toestemming te geven voor het verzamelen, gebruiken en bewaren van ‘bijzondere’ gegevens. Als u niet akkoord gaat, kunt u niet deelnemen aan dit onderzoek. Uw gegevens worden geanonimiseerd zodra de gegevensverzameling is voltooid. Dit doen wij door verzamelde gegevens die naar uw persoon te herleiden zijn (in het bijzonder uw IP-adres) direct te verwijderen uit de data. De geanonimiseerde gegevens worden minimaal 10 jaar bewaard op versleutelde servers van de onderzoeksuniversiteit: de Radboud Universiteit in Nederland. Vanwege het belang van controle, hergebruik en/of replicatie van onderzoeksresultaten worden onderzoeksgegevens (inclusief eventuele anonieme persoonsgegevens) steeds vaker gedeeld met of beschikbaar gesteld aan andere onderzoekers. Alleen anonieme gegevens vallen onder deze vorm van delen. Als u niet wilt dat uw geanonimiseerde gegevens worden gedeeld, kunt u niet deelnemen aan het onderzoek. Sommige personen en organisaties moeten toegang hebben tot uw persoons- en onderzoeksgegevens. Dit is nodig om te toetsen of het onderzoek goed en betrouwbaar is uitgevoerd. Deze personen en toezichthouders die uw gegevens ter verificatie inkijken zijn onder meer: bevoegde personen binnen de Radboud Universiteit (bijvoorbeeld een decaan, directeur of data officer) en (inter)nationale toezichthouders (bijvoorbeeld de Autoriteit Persoonsgegevens en het College voor Wetenschappelijke Integriteit). Zij zijn verplicht uw gegevens strikt vertrouwelijk in te zien. Voor deze toegang wordt u om toestemming gevraagd. Als u dit weigert, kunt u niet

deelnemen aan het onderzoek.

AVG

De Radboud Universiteit is verantwoordelijk voor de naleving van de Algemene Verordening Gegevensbescherming (AVG) bij de verwerking van jouw persoonsgegevens. De onderzoeker draagt er zorg voor dat jouw privacy en de daaraan verbonden voorwaarden worden gewaarborgd en houdt zich bij het uitvoeren van dit onderzoek aan de Nederlandse gedragscode wetenschappelijke integriteit en het universitaire beleid ten aanzien van de opslag en het beheer van persoons- en onderzoeksgegevens. De privacyverklaring van de Radboud Universiteit vind je op: <https://www.ru.nl/vaste-onderdelen/privacyverklaring-radboud-universiteit>. Bij vragen over jouw privacy kun je de Privacy Officer van de Faculteit der Sociale Wetenschappen benaderen (P.Janssen@socsci.ru.nl). Voor algemene vragen over dit onderwerp kun je contact opnemen met de Data Protection Officer van de Radboud Universiteit via privacy@ru.nl. Meer informatie over jouw rechten bij het verwerken van persoonlijke data kun je vinden op <https://www.ru.nl/privacy/english/protection-personal-data/data-subjects-rights/> en op de website van de Nederlandse Autoriteit Persoonsgegevens (<https://autoriteitpersoonsgegevens.nl>).

Het onderzoek is beoordeeld op basis van de Light Track Checklist van de Ethische Commissie Sociale Wetenschappen (ECSS) van de Radboud Universiteit. Op basis daarvan is er geen formeel bezwaar tegen dit onderzoek gebleken.

Jouw deelname

Jouw deelname aan dit onderzoek is geheel vrijwillig. Als je besluit niet deel te nemen, heeft dit geen gevolgen. Als je tijdens het onderzoek je toestemming wilt intrekken en je deelname wilt beëindigen, kun je dit doen door simpelweg het enquêtetabblad in je browser te sluiten. Eventuele resterende gegevens worden verwijderd. Als je vragen of opmerkingen hebt over dit onderzoek, stuur dan een e-mail naar: juul.vandekamp@ru.nl (onderzoeker) of paul.ketelaar@ru.nl (supervisor).

Toestemmingsverklaring

Ik bevestig dat:

- ik tussen de 16 en 29 jaar oud ben
- ik naar tevredenheid over het onderzoek geïnformeerd ben;
- ik de informatie goed heb gelezen;
- ik in de gelegenheid ben gesteld om vragen over het onderzoek te stellen;
- mijn eventuele vragen naar tevredenheid zijn beantwoord;
- ik goed over deelname aan het onderzoek heb kunnen nadenken;
- ik uit vrije wil deelneem aan het onderzoek.

Ik begrijp dat:

- ik het recht heb om mijn toestemming op ieder moment weer in te trekken zonder opgave van redenen en zonder dat dit nadelige gevolgen voor mij heeft, door het enquêtetabblad in de browser af te sluiten;
- mijn persoonsgegevens worden verwerkt volgens de AVG;
- mijn persoonsgegevens worden verwerkt volgens de privacyverklaring van de Radboud Universiteit (<https://www.ru.nl/vaste-onderdelen/privacyverklaring-radbouduniversiteit>);

Ik stem in dat:

- mijn persoons- en/of onderzoeksgegevens binnen dit onderzoek voor wetenschappelijke doelen worden verkregen en beschikbaar zullen zijn voor controle, hergebruik en replicatie;
- bijzondere persoonsgegevens over mij worden verwerkt;
- voor de controle van het onderzoek toezichthoudende autoriteiten mijn persoons- en onderzoeksgegevens kunnen inzien.

Ik stem in met deelname aan de studie

- ☐ Ja
- ☐ Nee

Bijlage B – Vragenlijst

1. Wat is uw leeftijd in cijfers?

○ ____

2. Wat is uw geslacht?

- Man
- Vrouw
- Anders

3. Wat is uw hoogst genoten opleiding?

- Primair onderwijs (basisschool en speciaal onderwijs)
- vmbo klas 1-4, havo/vwo klas 1-3, vavo (vmbo-niveau), mbo entree-opleiding, praktijkonderwijs
- havo/ vwo klas 4 – 6, vavo (havo/vwo-niveau), mbo (niveau 2-4)
- Associate degree (2-3 jarig hbo)
- wo-bachelor, hbo-bachelor, post-hbo opleiding
- wo-master, hbo-master
- Doctoraat (PhD)

4. Scenario

Beeld je in: je bent online aan het browsen en komt een advertentie tegen van een huidverzorgingsmerk. In dit onderzoek zie je één advertentie voor een crème van dit merk. Deze advertentie is ontwikkeld met behulp van AI.

Wij vragen je om de advertentie rustig en aandachtig te bekijken, alsof je deze in een echte online omgeving tegenkomt. Probeer tijdens het bekijken een indruk te vormen van zowel de advertentie als het merk achter het product. Na het bekijken van de advertentie volgt een korte vragenlijst. Neem dus rustig de tijd om de advertentie goed te bekijken. Je kunt na 5 seconden doorgaan naar de volgende pagina.



*Conditie met functioneel appeal.



*Conditioes met emotioneel appeal.

5. Meten variabele merkbetrouwbaarheid

- Met merk [x] krijg ik wat ik zoek in een crème
- [X] is een merk dat aan mijn verwachtingen voldoet.
- Ik heb vertrouwen in merk [x].
- [x] is een merk dat mij niet teleurstelt.
- Merk [x] zou eerlijk en oprecht omgaan met mijn zorgen.
- Merk [x] doet er alles aan om mij tevreden te stellen.
- Ik zou op merk [x] kunnen vertrouwen om een probleem op te lossen.
- Merk [x] vindt mijn tevredenheid belangrijk.
- Merk [x] zou mij op de een of andere manier compenseren voor een probleem met de crème.
- Merk [x] zou niet bereid zijn een probleem op te lossen dat ik met de crème zou kunnen hebben.

*5-punt Likertschaal van helemaal mee oneens tot helemaal mee eens

6. Meten variabele houding tegenover AI

- Ik sta positief tegenover toepassingen van AI.
- Ik vertrouw AI in het dagelijks leven.
- Ik ben sceptisch over AI.

*5-punt Likertschaal van helemaal mee oneens tot helemaal mee eens

7. Meten variabele betrokkenheid Huidverzorgingsproducten

- Ik ben geïnteresseerd in huidverzorgingsproducten.
- Huidverzorgingsproducten zijn belangrijk voor mij.

*5-punt Likertschaal van helemaal mee oneens tot helemaal mee eens

8. Einde onderzoek

Bedankt voor uw deelname aan dit onderzoek. Dit onderzoek richt zich op de manier waarop consumenten AI-gegenereerde reclame beoordelen en op de vraag onder welke omstandigheden deze vorm van reclame effectief kan zijn. Het wordt zeer gewaardeerd dat u de tijd heeft genomen om aan dit onderzoek deel te nemen.

Mocht u vragen hebben of meer informatie willen over het onderzoek, dan kunt u contact opnemen via juul.vandekamp@ru.nl