



Part. 1

미래를 보여주면 더 참여할까? 생성형 AI와 인간의 기부 예측 광고 효과 비교



책임 연구자 **송수진** 고려대학교





책임 연구자

송 수 진

아름다운재단 기부문화연구소 연구위원
고려대학교 글로벌비즈니스대학 교수

학력

미국 로드아일랜드 주립대학, 경영학 박사
미국 시몬스 컬리지, M.B.A.
KDI 국제정책대학원, 석사
고려대학교 정치외교학과, 학사

주요경력

한국경영학회 이사
한국마케팅학회 이사
한국마케팅과학회 이사
한국소비문화학회 상임이사
한국광고학회 이사
서울시 브랜드 자문위원
경제인문사회연구회 자문위원
과학기술정보통신부 민간위원
세종시 여성기업지원위원회 자문위원
LG 인화원 자문교수, MVP 과정 강사
동아비즈니스리뷰 (DBR) 객원편집위원, 필진
브랜드 매니저(Assistant Brand Manager), 마케팅 부서, 한국 P&G

저서

《소비자의 마음을 읽어드립니다》, e비즈북스(2023).

연구실적

- CSES 사회적 가치 연구공모전 논문상 최종 수상.
- 마케팅과학연구(Journal of Global Scholars of Marketing Science), 소비문화연구(Consumer Culture Research) 최우수 논문상.

-
- 고려대 미래연구기금(KU Future Research Grant), 융합연구기금(Interdisciplinary Research Grant) 수상.
 - 장애인이 포함된 광고는 더 멋질까: 포용적 마케팅의 브랜드 멋짐 효과 **마케팅연구**, (2023)
 - “Do female CMOs enhance firm performance? Power matters.” **Journal of Business Research** (2023)
 - “What Explains Smartwatch Adoption? A comparative study of South Korea and Indonesia”, **Asia Marketing Journal** (2022)
 - “Extending Diderot Unities: How Cosmetic Surgery Changes Consumption?”, **Psychology & Marketing** (2021)
 - “Motivation and Outcomes of Private Supplementary Tutoring”, **Korea Observer** (2020)
 - “A Study on The Motivation of Cyber Money Consumption of Interactive Media”, **Consumer Culture Research** (2018)
 - “When Women Are Dissatisfied: Gender Difference in the Effects of Failure Locus of Causality and Severity”, **Social Behavior and Personality** (2017)
 - “A Study on the Effects of Work addiction and Materialism on Brand Dependence”, **Consumer Culture Research** (2017)
 - “CEO Compensation and Concurrent Executive Employment of Outside Directors”, **KDI Journal of Economic Policy** (2016)
 - “Effects of Product Failure Severity and Locus of Causality on Consumers’ Brand Evaluation”, **Social Behavior and Personality** (2016)
 - “The Influence of the Philosophy of Science on Brand Loyalty”, **Journal of Global Scholars of Marketing Science** (2015)
 - “Higher Quality or Lower Price? How Value-Increasing Promotions Affect Store Reputation via Perceived Value”, **Journal of Business Research** (2014)
 - “When Disturbing Is Likable: Product Placement Effects on Multitasking Consumers,” **Journal of Advertising** (2011)
-

보조 연구자

송 채 원

고려대학교 일반대학원 기업경영학과 마케팅전공 박사과정생

요약

생성형 AI를 활용한 광고는 기업들에게 많은 이점을 제공한다. 특히 비영리 단체는 생성형 AI를 통해 시간과 비용을 절감할 수 있다. 또한 재난 상황과 같은 어려운 환경을 시각적으로 연출해야 하는 기부 광고에서, AI를 통해 위험한 현장을 찾아가거나 맞닥뜨리지 않아도 실재감 있는 광고를 생성할 수 있다. 뿐만 아니라 AI는 데이터를 분석하여 각 기부자의 관심사와 성향을 파악할 수 있어, 기부자의 마음을 끌 수 있는 맞춤형 메시지를 제공할 수 있다.

한편 소비자들은 생성형 AI가 만든 콘텐츠를 진정성이 부족하다고 인식하여 부정적으로 평가하는 경향이 있다. AI가 객관적인 작업을 수행하는 경우에는 비교적 높은 신뢰도를 보이지만, 창의성이나 감정이 개입되는 주관적인 작업을 수행하는 경우는 소비자들이 부정적으로 반응하기도 한다. 이를 알고리즘 기피 현상(algorithm aversion)이라고 한다.

이런 알고리즘 기피 현상에도 불구하고 생성형 AI는 기술의 고도화와 사용의 간편함, 업무 수행을 위한 학습 비용 및 실행 비용의 절감 효과 등으로 인해 지속하여 그 사용 범위와 깊이가 확장될 것으로 보인다. 따라서 본 연구에서는 소비자들이 보여온 알고리즘 기피 현상이 어떤 때에 감소하는지, 어떤 전략을 활용할 때 축소시킬 수 있는지 라는 질문을 가지고 실험을 진행하였다. 구체적으로, 고정 관념과 머신 휴리스틱 이론에 기반하여 AI가 예측 작업을 수행한다면, 소비자들은 이를 긍정적으로 인식할 것으로 판단했다. 이는 기부 광고에 있어서는 특히 정확성이 높은 전략이다. 기부 캠페인은 해당 기부를 통해 수혜자의 삶에 어떤 변화가 있겠는지, 기부 목적을 얼마만큼 달성할 수 있을지를 보여주는 미래 예측 방식을 자주 활용해왔다.

총 네 번의 실험으로 다음과 같은 결과를 얻었다. 첫째, 일반적인 기부 광고에서는 인간이 제작한 광고가 AI가 제작한 광고보다 신뢰도와 기부 의도가 높았다. 둘째, 광고에서 기부 결과를 예측해 보여주는 경우 제작자 유형(AI vs. 인간)에 따른 신뢰도와 기부 의도의 유의미한 차이는 없었으며, AI에 대한 신뢰도와 기부 의도가 평균적으로 더 높아져 알고리즘 기피 현상이 감소됨을 확인할 수 있었다. 셋째, 예측을 구체적으로 제시하는 경우, AI 제작자에 대한 신뢰도와 기부 의도가 인간 제작자보다 높아졌다. 그러나 예측이 추상적인 경우에는 AI 제작자와 인간 제작자가 만든 광고 사이의 차이

가 악화되었다.

본 연구는 알고리즘 기피 현상이 나타나기 쉬운 광고 제작 작업에서 이를 완화할 수 있는 전략을 제시함으로써 이론적 · 실무적 공헌점을 제공한다. 특히 AI가 예측 광고를 제작할 때 알고리즘 기피 현상이 완화된다는 점을 밝힘으로써, 향후 생성형 AI를 활용한 기부 캠페인 전략에 실무적 함의를 제시할 수 있을 것으로 기대한다.

Keyword: #생성형AI, #기부캠페인, #예측광고, #알고리즘기피, #알고리즘선호, #기부참여

1. 서론

최근 많은 기업들이 생성형 AI(Generative AI)를 활용하여 광고를 제작하고 있다. 하지만 소비자들은 광고와 같이 주관성이 개입되는 작업에 AI가 활용되는 것을 객관적인 업무에 AI를 활용할 때보다 부정적으로 받아들인다. 이러한 현상을 알고리즘 기피(algorithm aversion) 현상이라고 부른다(Castelo et al., 2019). 그럼에도 기업이 광고 제작에 생성형 AI를 활용하는 이유는 무엇일까? 생성형 AI를 활용하면 시간적 비용뿐만 아니라 금전적 비용도 줄일 수 있어 효율적으로 광고를 제작할 수 있다(Shah et al., 2020). 따라서 선행 연구는 생성형 AI가 만든 광고도 소비자에게 긍정적으로 인식될 수 있는 전략과 상황에 관한 연구를 진행해 왔다(예, Wu & Wen, 2021; Campbell et al., 2022; Brüns & Meißner, 2024).

기업들에 비해 비영리 단체나 공공 기관이 생성형 AI를 활용하는 경우는 아직까지는 많지 않다. 하지만 생성형 AI를 통한 광고 생성이 비용 절감의 효과가 있어 일반적으로 광고 예산이 제한된 비영리 단체가 사용할 경우에 효율성이 더 높을 것이다(Arango et al., 2023). 또한, 비영리 단체와 공공 기관이 자연재해나 전쟁과 같은 천재지변이나 재난 구호 사업 등을 진행한다는 점을 고려할 때, 직접 촬영이 어려운 지역과 상황을 설명하기 위해 AI로 생성한 이미지를 활용한다면 기부의 필요성을 광고를 통해 효과적으로 전달할 수 있다(Arango et al., 2023). 따라서 향후 비영리 기관 및 공공 기관에서 생성형 AI를 활용한 광고와 커뮤니케이션이 빈번히 사용될 수 있을 것이다.

소비행동학과 광고학 분야에서는 기부 행동과 AI를 연관지은 연구들이 최근 주목받고 있다. 챗봇과 인간을 비교하여 소비자의 기부 의도를 연구한 Zhou et al. (2022)는 챗봇을 활용할 경우 기부에 부정적인 영향을 미치는 것으로 보고했다. Lv & Huang (2024)은 AI가 구현해 낼 수 있는 특성인 개인화 서비스를 기부에 적용시켜 연구하였는데, 개인화 추천 서비스를 제공하는 경우에는 되려 기부 의도가 부정적으로 나타났다. 이는 기부를 어디에 할지 AI가 개인에게 추천해 주면 기부자들의 자율성이 손상을 입기 때문인 것으로 나타났다. 이것은 상업적 환경에서의 AI가 활용되는 것과 차이를 보이는 지점이다. 기업의 제품과 서비스를 사용하는 고객들은 자신에게 필요한 것이 무엇인지 상대가 분석하고 나에게 꼭 맞는 것을 추천해 주는 것을 선호한다. 즉 개인화 추천 서비스를 더욱 가치있게

여긴다. 그러나 기부와 같이 자신의 가치관과 생각을 투영한 항목(사람, 이슈, 사건 등)이 중요한 영역에서는 소비자들은 AI가 추천해 주는 것을 아직까지 신뢰하지 않는 것으로 보인다. Arango et al. (2023)은 공감 능력이 없는 AI가 기부 콘텐츠를 제작한 것에 대해 소비자들이 부정적으로 인식하고, 그에 따라 기부 의도 또한 낮아진다고 보고했다. 즉, 기부 캠페인에 AI가 개입되면 기부자들의 알고리즘 기피 현상이 나타나는 것이다.

본 연구는 위와 같이 선행 연구가 보고한 알고리즘 기피 현상이 감소하는 경우를 탐색하고자 한다. 계속 기술이 발전하고 소비자들 역시 사회 각 분야에서 생성형 인공지능의 사용과 활용을 경험하다 보면 기부 영역에서 생성형 인공지능이 활용된 광고에 대한 저항도 감소할 가능성이 있다. 그러나 이런 낯선 기술에 대한 자연스러운 저항 감소 이면에 소비자들은 근본적으로 AI와 인간이 만든 창작물을 어떻게 다르게 인식하고 판단하는지 그 심리와 행동 방식을 탐구할 필요는 여전히 존재한다. 되려, 생성형 인공지능의 활용이 보편화될 경우, 어떤 방식으로 활용하는 것이 효과적이고 어떤 경우에 소비자들의 부정적 반응을 야기할지에 대해 전략적 탐구가 선행될 필요가 있다.

머신 휴리스틱 이론에 따르면, 사람들의 AI에 대한 인식과 고정관념이 AI를 대하는 태도에 영향을 끼친다(Sundar & Kim 2019). 사람들은 AI가 객관적이고 정확성이 높다고 믿고 이런 특성이 나타날 때 AI를 긍정적으로 인식한다. 그러나, AI는 감정 지능이 낮다고 인식한다(Holthöwer & van Doorn, 2023). 따라서, AI가 예측과 같이 객관적인 업무를 수행할 때는 신뢰하고, 창작과 같은 주관적인 업무를 수행할 때는 알고리즘 기피 현상을 보인다(Castelo et al., 2019). 이런 선행 연구와 추론을 바탕으로, 본 연구는 예측이 필요한 광고, 예측이 삽입된 기부 광고에 대해서는 소비자들 AI가 만든 광고를 더욱 신뢰하고 기부의 또한 높을 것으로 예상했다. AI의 뛰어난 계산 능력으로 인해서(Kim et al., 2021) 예측과 관련한 많은 분야에서는 알고리즘이 인간보다 더 정확한 것으로 인식할 것이기 때문이다(Grove et al., 2000; Jung & Seiter 2021). 이런 추론과 선행 연구를 바탕으로 본 연구는 AI가 사용된 기부 광고에서 알고리즘 기피 현상이 감소시키는 전략을 밝혀 이론적, 실무적 시사점을 도출하려 한다.

2. 이론적 배경

1) 기부 광고에서 AI의 활용

비영리 단체가 AI를 활용할 경우 많은 장점이 있다. 첫째, AI는 대규모 데이터를 분석하여 기부 가능성이 높은 타겟 그룹을 식별하고, 이들을 대상으로 효율적인 광고와 메시지를 전달할 수 있다. 둘째, 비교적 광고 예산이 적은 비영리 단체는 AI를 활용해 광고 제작과 집행 비용을 절감할 수 있다(Arango et al., 2023). 셋째, 자연재해나 전쟁과 같은 위험 상황을 묘사할 때 AI를 활용하면 촬영에 따른 위험과 금전적 비용을 줄이면서도 현실감 있게 표현할 수 있다. 넷째, 광고 제작뿐만 아니라 광고 캠페인의 효과를 모니터링하여 대응할 수 있다. 다섯째, 기부자 데이터 분석을 통해 개인화된 광고나 메시지를 삽입하는 등 개인화된 콘텐츠를 제공하고, 상호작용을 증진하며, 기부 경험을 향상할 수 있다(Gao & Liu, 2023). 마지막으로 챗봇을 통해 인건비를 절감하고, 시간 제한 없는 24시간 상담 서비스를 제공할 수 있다.

그러나 기부 맥락에서 AI의 활용은 다소 부정적인 반응을 불러일으키기도 한다. 예를 들어, AI가 상품이나 서비스에 대한 개인화 추천을 제공하는 경우 소비자들은 이를 자신과 높은 관련성이 있다고 느끼며, 이는 브랜드 태도에 긍정적인 영향을 미친다(De Groot, 2022). 그러나 기부 추천 서비스에서는 기부자들이 자율성을 침해받는다고 느껴 부정적인 인식을 형성하는 경우가 있다(Lv & Huang, 2024). 또한 기부자가 챗봇과 상호작용할 때는 인간과 상호작용할 때보다 기부 의도가 감소하는 경향이 있는데, 이는 챗봇이 기부자에 대한 도덕적 판단을 하는 정도가 낮기 때문이다(Zhou et al., 2022). 마지막으로, AI가 제작한 기부 광고는 인간이 제작한 기부 광고보다 감정이입(empathy)을 덜 불러일으켜 기부 의도가 낮은 것으로 나타났다(Arango et al., 2023). 이처럼 비영리 단체가 AI를 활용하는 것은 많은 장점이 있지만, 아직까지는 기부자들이 부정적인 인식을 갖고 있는 경향도 보고되고 있기에 맥락과 상황에 맞는 판단과 적절한 어조를 구분하여 섬세하게 접근하는 등 주의가 요구된다.

2) 정보원의 신뢰(Source credibility)

정보원의 특성은 수신자의 설득력 판단에 중요한 영향을 미치며(Ismagilova et al., 2020), 특히 정보원에 대한 신뢰도가 높을수록 소비자는 해당 광고를 더욱 신뢰하는 경향이 있다(Chaiken & Maheswaran, 1980). 귀인 이론에 따르면 소비자는 외부 정보에 대한 정확한 인식을 얻으려는 동기를 가지며(Kelly, 1972), 신뢰도가 높은 메시지들을 위주로 처리하려는 경향이 있다(Gotlieb & Sarel, 1991). 신뢰도가 높은 정보원은 신뢰도가 낮은 정보원보다 더 큰 설득력을 가지는 것으로 알려져 있다(Hovland et al., 1953).

최근 SNS의 확산으로 광고 제작자는 물론 인플루언서나 일반 사용자들도 사용자 생성 콘텐츠(UGC)를 제작하게 되면서, 정보원의 신뢰도 판별이 중요해졌다(Lee et al., 2017). 인플루언서의 광고 콘텐츠가 높은 신뢰도를 보일 경우, 그 인플루언서가 홍보하는 제품에 대한 구매 의도가 증가한다(Lou & Yuan, 2019). 가상 인플루언서와 인간 인플루언서를 비교했을 때, 인간 인플루언서가 가상 인플루언서보다 더 높은 신뢰도를 얻었으며, 가상 인플루언서가 팔로워들과 상호작용성이 높을 때 신뢰도 역시 상승하는 것으로 보고되었다(Yang et al., 2023). 소비자가 제작하는 UGC의 경우에는 UGC의 신뢰도가 높을수록 구매 의도와 온라인 구전 의도도 증가한다(Muda & Hamzah, 2021). 또한, 동영상의 품질이 낮은 경우에 소비자들은 사용자 생성 콘텐츠(UGC)에 대해 에이전시가 제작한 광고보다 더 높은 신뢰를 보였다(Hautz et al., 2014).

데이터의 양이 증가함에 따라 미디어 회사들은 기사 작성 또는 소셜 미디어에 콘텐츠를 게시하는데 알고리즘과 인공지능(AI)을 이전보다 더 많이 활용하고 있다(Hofeditz et al., 2021). 콘텐츠를 창작하고 제작하는 일은 더 이상 인간만의 영역이 아니고, AI도 중요한 역할을 할 수 있는 분야가 되었다. 하지만 AI가 가짜 뉴스를 생성할 수도 있기 때문에 생성된 콘텐츠의 신뢰성을 식별하는 것이 점점 중요해지고 있다(Chiang et al., 2022). 특히 실제 사건을 매우 정교하게 모방하여 진위를 구별하기 어려운 딥페이크(‘딥러닝’과 ‘가짜’의 합성어) 콘텐츠가 증가함에 따라 사람들은 콘텐츠의 진위 여부를 구별하는 데 어려움을 겪고 있다(Whyte, 2020). 이러한 이유로 생성형 AI가 만든 콘텐츠에 대한 신뢰도는 생성형 AI가 제작한 콘텐츠에 대한 수용 여부를 판별하는 데 중요한 요인이 된다.

3) AI에 대한 소비자 인식

기업들은 생성형 AI를 다양한 업무에 활용하고 있다. 대표적인 예로, 엘지 유플러스(LGU+)는 시나리오 작성부터 영상 생성 및 편집까지 생성형 AI 기술로 만든 광고 영상을 공개했다. 이 광고는 마케터 역할을 맡은 배우 주현영이 AI에게 광고 제작을 맡기는 줄거리로, 기존 방식에 비해 제작 기간이 3분의 1로 단축되었다(임지선, 2023). 또 다른 사례로, 배스킨라빈스는 구글의 생성형 AI 제미니(Gemini)를 통해 새로운 맛의 아이스크림을 개발했다. 아마존은 생성형 AI를 개인화 추천 서비스에 도입하여 실시간으로 고객의 상황과 맥락에 맞는 맞춤형 추천 서비스를 제공하고 있다.

고정 관념은 평가 대상이 특정 집단에 속한다고 여겨질 때, 그 집단에 대한 일반적인 생각이나 믿음을 바탕으로 다른 사람을 빠르게 판단하는 과정이다(Im & Lee, 2023). 고정 관념 이론은 사람뿐만 아니라 브랜드(Kolbl et al., 2020), 국가(Khachaturian & Morganosky, 1990), 컴퓨터(Nass et al., 1997) 또는 AI와 같은 비인간 존재에 대한 추론과 인지 과정을 설명하는 데도 적용된다(Ahn et al., 2022).

머신 휴리스틱(machine heuristic) 이론에 따르면, 기계에 대한 인식이나 고정 관념은 기계가 수행하는 작업물에 대한 기대에 영향을 미친다(Sundar & Kim, 2019). 사람들은 기계에 대해 객관적이고 정확하다고 인식하는 경향이 있다(Sundar & Kim 2019). 이러한 이유로 정량화되고 측정 가능한 사실이 포함된 객관적인 업무를 수행하는 경우에는 알고리즘에 대한 신뢰도가 높다고 인식하지만, 감정이나 직관이 필요한 주관적인 업무를 수행하는 경우에는 신뢰도가 낮다고 인식하고, 인간이 수행하는 것을 더 선호한다(Castelo et al., 2019). 이렇게 AI를 신뢰하지 못하는 현상을 알고리즘 기피(algorithm aversion)라고 한다(Castelo et al. 2019).

최근 AI는 광고의 효과를 데이터로 측정하고 분석하는 역할을 수행함과 동시에 창의성이 요구되는 광고 제작 주체로도 활용되고 있다(Qin & Jiang, 2019). 생성형 AI의 획기적인 발전으로, 음원과 영상, 이미지와 콘텐츠를 기획 개발하는 광고 제작 업무에서 AI의 사용률과 중요도는 계속 증가할 것이다(Wu et al., 2024). 비용과 시간을 절감할 수 있는 장점이 있지만, 소비자들은 AI가 제작한 광고에 대해 부정적인 인식을 갖는다(Shah et al., 2020). 소비자들은 광고 제작을 창의성이 요구되는 업무로, 주관적인 요소가 많은 작업으로 인식하고 있기 때문이다(Wu et al., 2024). 따라서, AI가 콘텐츠

츠 제작에 개입했다고 알려지면 콘텐츠의 진정성이 낮아지고, 브랜드에 대한 태도 역시 부정적으로 형성된다(Brüns & Meißner 2024). 기부 광고의 경우, AI가 제작한 광고는 소비자들의 기부 상황과 주제에 대한 공감에 부정적인 영향을 미쳐 기부 의도에도 부정적 영향을 미칠 수 있다(Arango et al., 2023).

하지만 앞서 언급하였듯이 AI가 수행하는 모든 작업에서 알고리즘 기피 현상이 발생하는 것은 아니다. AI의 정확성이 높게 인식되는 분야, 예를 들어 숫자 추정이나 예측 업무에서 사람들은 인간보다 AI를 신뢰하고 호감을 갖는다(Castelo et al. 2019, Logg et al. 2019). 또한 사람들은 AI가 객관적으로 업무를 수행한다고 인식하기 때문에 인간보다 성차별이나 인종 차별 같은 편향된 의사 결정을 내리지 않을 것으로 믿는다(Bonezzi & Ostinelli, 2021). Logg et al. (2019)은 이런 소비자들의 심리와 태도를 알고리즘 기피(algorithm aversion)와 반대로 알고리즘에 대해서 긍정적인 감정을 느끼게 되는 경우로 ‘알고리즘 존중(algorithm appreciation)’으로 명명하였다.

종합적으로, 선행 연구 결과는 AI에 대한 인식이 작업 유형에 따라 긍정적이거나 부정적으로 달라질 수 있음을 시사한다. 광고 제작 시 AI가 개입되었을 때, 알고리즘 기피 현상이 나타날 것이다. 그러나 AI가 인간보다 유능하다고 여겨지는 예측(Castelo et al., 2019; Logg et al., 2019)이 들어간 광고를 AI가 제작한 경우는 오히려 알고리즘 기피 현상이 감소할 것이다.

H1. 기부 광고를 AI가 제작하는 경우는 인간이 제작하는 경우에 비해, (a) 신뢰도와 (b) 기부 의도가 낮을 것이다.

H2. 예측 광고를 제작할 경우, 신뢰도와 기부 의도에 있어 광고 제작자 유형(AI vs. 인간) 간의 차이가 약화될 것이다. 즉, 광고 제작자 유형에 관계없이 (a)신뢰도와 (b)기부 의도에 유의미한 차이가 없을 것이다.

H3. 광고 유형과 광고 제작자 유형의 상호작용이 기부 의도에 미치는 영향은 신뢰도에 의해 매개될 것이다. 즉, 광고 유형과 광고 제작자 유형 간의 상호작용은 신뢰도에 긍정적인 영향을 미치며, 신뢰도는 기부 의도에 긍정적인 영향을 미칠 것이다.

3. 실험 1

1) 실험 설계

알고리즘 기피 현상은 AI가 제작한 기부 광고에 대해 부정적인 인식을 초래하지만, AI의 예측 능력과 같은 강점을 광고에 명시적으로 강조함으로써 이러한 부정적 인식을 완화시킬 수 있을 것으로 예상된다. 따라서 위 가설들을 검증하기 위해서 ‘기부세상’이라는 가상의 비영리 단체를 만들었으며, 시나리오를 통해서 조작하였다. AI가 광고를 제작하는 경우에는 ‘이 광고는 AI가 제작하였습니다’라는 문구를 기재하였다. 예측 광고가 아닌 경우에는 ‘안녕하세요, 비영리 단체 기부세상입니다. 저희는 기부자들의 기부금을 해외 보건 의료를 지원하는 데 사용하고 있습니다. 이 사업을 통해 굶주림과 영양실조로 고통받는 어린이들에게 필요한 영양식과 의료 서비스를 제공하고 있습니다. 앞으로 도움이 필요한 아이들을 위해 저희 기부세상의 활동에 많은 관심과 지지를 부탁드립니다.’라는 시나리오를 제공하였다. 예측 광고인 경우에는 ‘AI는(사람들은) 2025년에 가뭄과 홍수 같은 극단적 기상 이변으로 식량 부족이 심화될 것으로 예측합니다. 사람들은 극단적 기상 이변으로 도움이 필요한 어린이들의 수가 2024년 대비 2025년에 약 15% 증가할 것으로 예측합니다. ‘라는 시나리오를 추가로 제공하였다.

2) 절차 및 조작 점검

실험 1은 2(광고 제작자 유형: 인간 vs. AI)×2(광고 유형: 예측 vs. 비예측)으로 집단 간 설계로 구성하였다. 피실험자들은 각 네 가지 조건 중 한 조건에 무작위로 할당되었다. 실험 1에서 사용될 실험 적절성을 판단하기 위해 AI의 광고인 경우에는 ‘위 광고는 누가 제작하였습니까?’라는 질문을 통해서 조작 점검을 하였다. 그 결과, 2명의 참가자를 제외하고 광고 제작자를 실험 조건과 일치하게 응답하였다. 예측 광고의 경우에는 ‘이 광고에는 2025년을 예측한 광고 문구가 언급되어 있습니까?’

라는 질문을 통해 조작 점검을 실시하였으며, ‘아니오’ 라고 응답한 2명을 제외하고 예측 조건을 맞게 대답하였다. 이에 따라 광고 유형의 조작이 성공적으로 이루어졌음을 확인할 수 있었다. 본 연구에서 활용한 측정 변수 및 문항은 다음 <표 1>과 같다.

본 실험은 한국의 대학교 재학생들을 대상으로 조사하였다. 131명의 피실험자 중 조작적 점검 문항에 제대로 답변을 하지 않았거나 성실하게 답변하지 않은 5명의 참가자를 제외한 126명(여성: 63.2%, 20대: 96.0%)을 대상으로 분석을 진행하였다.

<표 1>

측정 변수	측정 문항	Cronbach's alpha	참고문헌
신뢰도	이 광고는 신뢰성을 가지고 있다/신뢰할 만하다.	.914	Cotte et al. (2005)
기부 의도	이 [브랜드 이름]에 기부할 의도가 있다. 이 [브랜드 이름]에 기부할 예정이다.	.945	Ye et al. (2015)

3) 연구 결과

신뢰도와 기부 의도에 대한 광고 제작자 유형(AI vs. 인간)과 광고 유형(예측 vs. 비예측) 각각의 주효과와 둘 간의 상호작용 효과를 검증하기 위해 이원 배치 분산분석을 실시하였다. 그 결과 신뢰도에 대한 광고 제작자 유형의 주효과($F(1,121) = .453, p = .502$)와 광고 유형의 주효과($F(1,121) = 3.425, p = .067$)는 유의미하지 않는 것으로 나타났다. 기부 의도에 대한 광고 제작자 유형의 주효과($F(1,121) = .226, p = .636$)와 광고 유형의 주효과($F(1,121) = 3.148, p = .079$)도 유의미하지 않는 것으로 나타났다. 반면에 광고 제작자 유형과 광고 유형의 상호작용 효과는 신뢰도($F(1,121) = 6.623, p = .011$)와 기부 의도에 유의미한 영향을 미치는 것으로 나타났다($F(1,121) = 4.420, p = .038$) (<그림 1> 참조).

상호작용 효과를 본페로니 다중 비교(Bonferroni's multiple comparison)를 통해 확인한 결과는 다음 <표 2>와 같다. 비예측 광고의 경우 AI($M=3.17$, $SD=.259$)보다 인간($M=4.03$, $SD=.263$)이 제작한 광고에 대한 신뢰도가 유의미하게 더 높게 나타났다($p=.021$). 반면, 예측 광고에서는 AI와 인간 간에 유의미한 차이가 없는 것으로 나타났다($p=.19$). 기부 의도의 경우, 인간이 제작한 비예측 광고가 AI가 제작한 경우보다 더 높은 기부 의도를 불러일으켰으나($M_{인간}=3.37$, $SD=.256$, $M_{AI}=2.70$, $SD=.26$), 유의 확률이 0.065로 통상 0.05수준에서 검증하는 수준에서의 유의미한 차이는 보이지 않았고 유의성에 근접한 결과를 확인하였다. 예측 광고에서는 인간과 AI 간 유의미한 차이가 없었다($p=.261$). 따라서 광고 제작자 유형(AI vs. 인간)에 따른 신뢰도와 기부 의도의 차이는 비예측 광고일 때 더 강하게 나타나며, 예측 광고에서는 이 차이가 약화되는 경향이 있음을 알 수 있다.

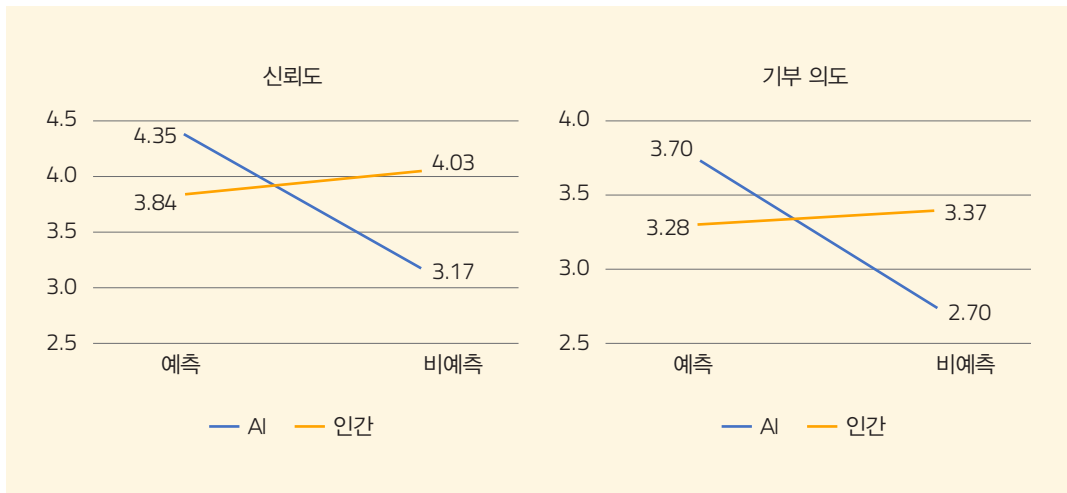
H3인 신뢰도의 매개 효과를 검증하기 위해 Hayes(2017)가 제안한 부트스트랩 방법(Process model 8)을 실시하였다. 광고 제작자 유형(X)을 독립 변인, 광고 유형(W)을 조절 변인으로 설정하여 검증한 결과, 상호작용이 기부 의도에 미치는 영향이 유의하지 않았다($Effect=.2771$, $S.E.=.4237$, $LLCI=-.5619$, $ULCI=1.1161$). 하지만 매개 변인인 신뢰감(M)이 추가된 모형에서 광고 제작자 유형과 광고 유형의 상호작용이 기부 의도에 미치는 간접 효과는 95%의 신뢰 구간 내에서 유의하게 나타나($Effect=.8121$, $S.E.=.3234$, $LLCI=.1875$, $ULCI=1.4541$) 신뢰감의 완전 매개 효과를 확인하였다. 이에 따라 H3도 지지되었다.

4) 논의

실험 1의 주요 결과는 다음과 같다. AI가 광고를 제작하는 경우 일반적으로 알고리즘 기피 현상이 발생하나, AI의 강점인 예측 능력을 광고에 명시적으로 강조했을 때, 이 현상이 감소하는 것을 확인할 수 있었다. AI가 제작한 예측 광고는 신뢰도가 향상되며, 기부 의도 또한 높아지는 것으로 나타났다. 또한 평균적으로 AI가 제작한 예측 광고에 대한 신뢰도와 기부 의도가 인간보다 더 높은 것으로 나타났다. 이는 AI가 특정 능력에서 인간보다 우수하다고 인식될 때, 그 기능이 강조된 광고를 통해 소비자의 알고리즘에 대한 부정적 인식을 줄이고 긍정적인 태도로 전환할 수 있음을 시사한다. 따

라서 AI가 제작한 광고에 AI의 예측 능력과 같은 강점을 명확히 표현하는 것이 알고리즘 기피 현상을 완화하며, 소비자들의 긍정적인 반응을 이끌어 낼 수 있는 효과적인 전략이 될 수 있다.

〈그림 1〉



〈표 2〉

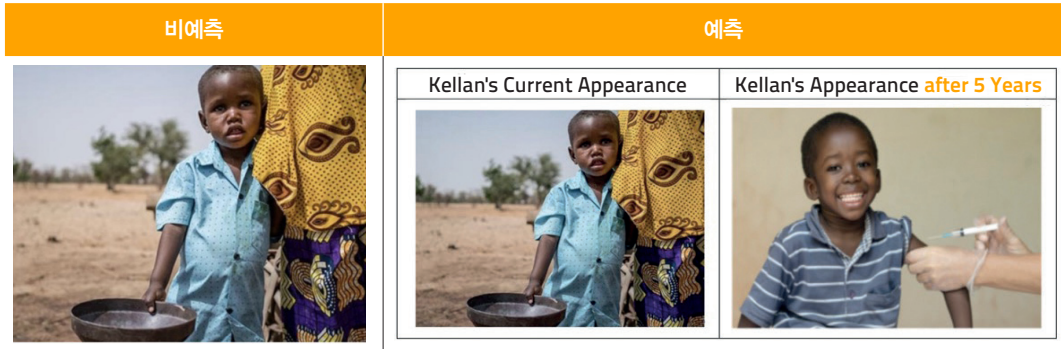
종속 변수	광고 유형	광고 제작자 유형	표본 수	평균	표준 오차	유의 확률
신뢰도	비예측	AI	31	3.17	.259	.021
		인간	32	4.03	.263	
	예측	AI	30	4.35	.276	.190
		인간	33	3.84	.267	
기부 의도	비예측	AI	31	2.70	.260	.065
		인간	30	3.37	.256	
	예측	AI	29	3.70	.269	.261
		인간	33	3.28	.252	

4. 실험 2

1) 실험 설계 및 절차

실험 2와 실험 3은 실험 1의 결과를 바탕으로 AI가 예측한 미래 상황을 텍스트가 아닌 사진으로 제시할 때 기부자들의 기부 의도가 어떻게 달라지는지를 살펴보려고 한다. 이는 사진 제작이 텍스트 작성보다 더 높은 창의성을 요구하는 작업으로 인식될 수 있어, 기부자들에게 더 큰 알고리즘 기피 현상을 유발할 가능성을 확인하고자 진행한 실험이다. 또한 실험 2에서는 인간의 미래 예측을, 실험 3에서는 환경의 미래 예측과 인간의 미래 예측 시나리오로 소비자들의 인식과 태도를 점검할 것이다. 이는 고정관념 이론과 머신 휴리스틱 이론을 현실에서의 소비자들의 AI 관련 경험과 인식에 적용하여 가설을 좀더 정교히 확인하기 위함이다. 환경을 예측하는 작업은 AI를 활용한 날씨 예측이 이미 보편적으로 사용되고 있어 AI가 수행하는 대표적인 예측 작업으로 여겨진다. 반면 인간의 미래를 예측하는 것은 통제하기 어려운 변수들이 많아 주관성이 강한 분야로 인식된다. 특히 소비자들의 인간의 미래에 대한 예측 경험은 타로나 사주, 점처럼 주관적이고 창의적인 스토리가 개입되는 것으로 인식하는 경우가 잦다. 따라서 주관성이 크게 개입되는 것으로 여겨지는 인간의 미래를 AI가 예측한 경우에도 실험 1에서와 비슷한 결과가 나타나는지 확인하고자 하였다. 실험 2에서 사용된 자극물은 다음 <그림 2>와 같다. 시나리오의 경우에는 비예측 조건인 경우에는 ‘아프리카에 살고 있는 Kellan을 도와주세요. Kellan은 열악한 보건 시스템으로 인해서 2019년부터 백신을 하나도 맞지 못하고 있습니다. Kellan이 백신을 맞을 수 있도록 도와주세요’라고 언급하였다. 예측 조건인 경우에는 ‘직원(AI)은 여러분의 기부로 앞으로 5년 동안 Kellan이 백신을 접종 받아 다양한 질병을 예방할 수 있을 것으로 예측하고 있습니다.’ 라는 시나리오를 추가적으로 제공하였으며, 두 조건 모두 AI인 경우에는 ‘AI가 제작하였습니다.’라고 언급하였다.

〈그림 2〉



실험 1과 마찬가지로 2(광고 제작자 유형: 인간 vs. AI)×2(광고 유형: 예측 vs. 비예측)으로 집단 간 설계로 구성하였다. 피실험자들은 각 네 가지 조건 중 한 조건에 무작위로 할당되었다. 실험 2에서 사용될 실험 적절성을 판단하기 위해 AI의 광고인 경우에는 ‘위 광고는 누가 제작하였습니까?’라는 질문을 통해서 조작 점검을 하였다. 70명 중 4명을 제외하고는 참가자들이 AI가 제작하였다고 응답하였으며, 광고 유형의 조작이 성공적으로 이루어졌음을 확인할 수 있었다. 측정 문항은 실험 1과 똑같이 신뢰도($\alpha = .952$)와 기부 의도($\alpha = .942$)를 측정하였다.

본 실험은 설문 기관 Prolific을 사용하여 참가자들을 모집하였다. 140명의 참가자가 참여하였지만, 이들 중 조작적 점검 문항에 제대로 답변을 하지 않았거나 성실하게 답변하지 않은 6명의 참가자가 제외되었다. 따라서 본 연구는 134명(남자: 55%, white: 55%, 나이: $M=33.27$ 세)을 대상으로 분석을 진행하였다.

2) 연구 결과

신뢰도와 기부 의도에 대한 광고 제작자 유형(AI vs. 인간)과 광고 유형(예측 vs. 비예측) 각각의 주효과와 둘 간의 상호작용 효과를 검증하기 위해 이원 배치 분산분석을 실시하였다. 신뢰도와 기부 의도에 대한 광고 제작자 유형의 주효과는 유의미하지 않는 것으로 나타났다(신뢰도: $F(1,130)=.963, p=.328$, 기부 의도: $F(1,130)=3.947, p=.597$). 반면에 광고 유형의 주효과는 유의미한 것으로 나타났다(신뢰도: $F(1,130)=7.147, p=.008$, 기부 의도: $F(1,130)=4.571, p=.034$)(〈그림 3〉 참조).

광고 제작자 유형과 광고 유형의 상호작용 효과는 신뢰도($F(1,130)=6.028, p=.015$)와 기부 의도($F(1,130)=6.821, p=.01$)에 유의미한 영향을 미치는 것으로 나타났다. 상호작용 효과가 어떻게 나타났는지 본페로니 다중 비교(Bonferroni's multiple comparison)를 통해 확인한 결과는 다음 〈표 3〉과 같다. 비예측 광고의 경우에는 AI($M=4.45, SD=.265$)보다 인간($M=5.36, SD=.249$)에 대한 신뢰도가 더 높은 것으로 나타났다($p=.017$). 하지만 예측 광고의 경우에는 두 집단 간의 유의미한 차이가 없는 것으로 나타났다($p=.296$). 비예측 광고에 대한 기부 의도의 경우에는 AI와 인간간의 유의미한 차이가 나타났으며($p=.031$), 인간($M=4.4, SD=.275$)이 AI($M=3.54, SD=.293$)보다 더 높은 것으로 나타났다. 반면에 예측 광고의 경우에는 광고 제작자 유형 간의 유의미한 차이가 없는 것으로 나타났다($p=.14$). 이에 따라 H1과 H2 모두 지지되었다.

신뢰도의 매개 효과를 검증하기 위해 Hayes(2017)가 제안한 부트스트랩 방법(Process model 7)을 실시하였다. 광고 제작자 유형(X)을 독립 변인, 광고 유형(W)을 조절변인으로 설정하여 검증한 결과, 상호작용이 기부 의도에 미치는 영향이 유의하지 않았다($Effect=.0341, S.E.=.2326, LLCI=-.4265, ULCI=.4948$). 하지만 매개 변인인 신뢰감(M)이 추가된 모형에서 광고 제작자 유형과 광고 유형의 상호작용이 기부 의도에 미치는 간접 효과는 95%의 신뢰 구간 내에서 유의하게 나타나($Effect=.9789, S.E.=.4322, LLCI=.1976, ULCI=1.8783$) 신뢰감의 완전 매개 효과를 확인하였다. 이에 따라 H3가 지지되었다.

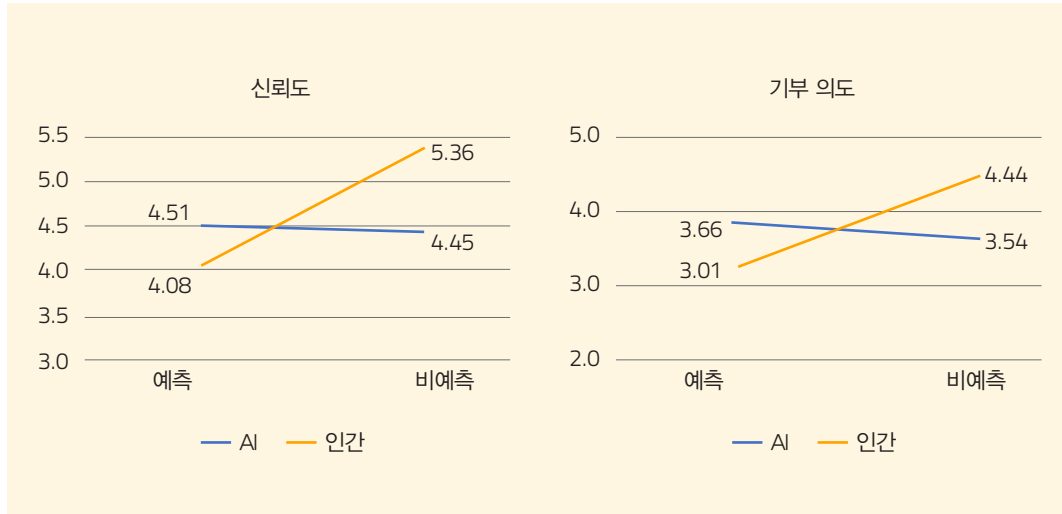
3) 논의

실험 2는 실험 1의 연구 설계를 바탕으로 광고 내용을 텍스트에서 사진으로 변경하여, AI가 시각적으로 미래를 예측하는 방식을 채택하였다. 연구 결과, 미래 예측을 보여주지 않는 광고에서는 인간이 제작한 광고와 AI가 제작한 광고 간에 유의미한 차이가 나타나, 알고리즘 기피 현상이 발생한 것으로 볼 수 있다. 그러나 미래 예측을 보여주는 광고의 경우 AI와 인간이 제작한 광고 간 유의미한 차이가 없었으며, 이는 실험 1과 동일하게 AI가 예측을 포함한 광고를 제작하는 경우 알고리즘 기피 현상이 감소함을 시사한다. 특히, 더 높은 창의성이 요구되는 사진 제작과 인간 예측을 포함한 경우에도 알고리즘 기피 현상이 줄어들었으며, 평균적으로 AI가 제작한 광고가 인간보다 더 높은 신뢰도와 기부 의도를 유발했음을 확인할 수 있었다.

〈표 3〉

종속 변수	광고 유형	광고 제작자 유형	표본 수	평균	표준 오차	유의 확률
신뢰도	예측	AI	32	4.51	.249	.296
		인간	34	4.08	.249	
	비예측	AI	34	4.45	.265	.017
		인간	34	5.36	.249	
기부 의도	예측	AI	32	3.66	.275	.140
		인간	34	3.01	.275	
	비예측	AI	30	3.54	.293	.031
		인간	34	4.44	.275	

〈그림 3〉



5. 실험 3

1) 가설

실험 3에서는 이전 실험들을 바탕으로 기부 예측의 대상과 예측 방식에 따라 AI와 인간이 제작한 광고에 대한 소비자들의 태도에 차이가 있는지 탐색하고자 한다. 실험 3A는 환경 미래 예측 시나리오로 실험을 진행하고 실험 3B는 인간 미래 예측 시나리오로 실험을 진행했다. Kim et al. (2021)은 AI의 추천에 대해 사람들은 더 구체적으로 표기된 숫자에 대해 높은 신뢰감을 보이고, 더 긍정적으로 평가한다고 보고한다. 실험 1과 실험 2에서는 AI가 예측 광고를 만든 경우에 인간이 예측 광고를 만든 경우와 비교하여 신뢰도와 기부 의도가 높음을 보였다. 이에 실험 3에서는 AI가 제작한 예측 광고

라 할지라도 기부 예측 대상과 예측 방식에 따라 인간 제작자와의 차이가 심화되거나 약화될 수 있을 지를 알아보고자 한다. 즉, 선행 연구와 실험 1, 2의 결과와 추론을 바탕으로 AI가 예측 광고를 구체적인 정보와 함께 제작한 경우에는 인간이 제작한 경우보다 더 높은 신뢰도와 기부 의도를 이끌어 낼 것이라 예상된다. 따라서 다음과 같은 가설을 설정하였다.

H4. 예측 광고가 구체적일수록 AI가 제작한 광고와 인간이 제작한 광고 간의 (a) 신뢰도와 (b) 기부의도의 차이가 클 것이다.

H5. 예측 광고가 추상적일 경우에는 AI 제작자와 인간 제작자 간의 (a) 신뢰도와 (b) 기부 의도의 차이가 적을 것이다.

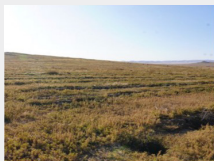
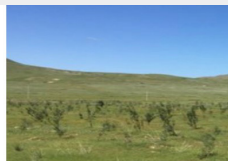
2) 실험 3A 설계 및 절차

실험 3A에서는 메시지 구체성에 따른 기후 변화 예측에 대한 소비자 태도를 살펴보고자 하였다. 따라서 실험 3은 2(광고 제작자 유형: AI vs. 인간)×3(광고 유형: 구체적인 예측 vs. 추상적인 예측 vs. 비예측)으로 집단 간 설계로 구성하였다. 실험 3A에서 사용된 자극물은 다음 <그림 4>와 같으며, AI인 경우에만 AI가 제작하였다고 명시하였다. 참가자들은 Prolific을 통해서 모집하였으며, 여섯 개의 조건 중 하나를 무작위로 할당되었다.

실험 3A에서 사용될 실험 적절성을 판단하기 위해 AI의 광고인 경우에는 ‘위 광고는 누가 제작하였습니까?’ 라는 질문을 통해서 조작 점검을 하였다. 105명 중 3명을 제외하고는 참가자들이 AI가 제작하였다고 응답하였으며, 광고 유형의 조작이 성공적으로 이루어졌음을 확인할 수 있었다. 측정 문항은 이전 실험과 동일하게 신뢰도($\alpha = .703$)와 기부 의도($\alpha = .779$)를 측정하였다.

본 실험은 220명의 참가자가 참여하였고, 조작적 점검 문항에 제대로 답변을 하지 않았거나 성실하게 답변하지 않은 18명의 참가자는 제외하였다. 본 연구 참가자의 인구통계학적 특성은 202명(여성(61.4%), 백인(90.6%), 나이($M=40.2$ 세))이다.

〈그림 4〉

	AI	인간
비예측	 The photo of Mongolia created by AI	 The photo of Mongolia taken in 2024
	Mongolia is facing a harsh reality: its landscapes are turning into desert, its ecosystems are crumbling, and its communities are being uprooted due to climate change. In response, Greentree is launching a vital tree-planting initiative across the nation. Join us in this crucial mission. We are not just planting trees; we're also bringing life back to the land and helping to heal our planet. Your donation can make a real difference. Help us transform Mongolia and restore hope to its people.	
구체적인 예측	 The photo of Mongolia's state in 2024  The photo of Mongolia's state in 2026, estimated by AI	 The photo of Mongolia's state in 2024  The photo of Mongolia's state in 2026, estimated by an employee
	Mongolia is facing a harsh reality: its landscapes are turning into desert, its ecosystems are crumbling, and its communities are being uprooted due to climate change. In response, Greentree is launching a vital tree-planting initiative across the nation. This right image depicts Mongolia in 2026, as predicted by AI(employee). According to the AI(employee), planting 20,000 trees could transform this desertification into grasslands. These trees are expected to absorb 160,000 kilograms of carbon dioxide, potentially lowering Mongolia's temperature by about 2 degrees Celsius. This change will create the conditions for approximately 500 species of animals and plants to thrive. Imagine how your help and donations can change the future of Mongolia. We are not just planting trees; we're also bringing life back to the land and helping to heal our planet. Your donation can make a real difference. Help us transform Mongolia and restore hope to its people.	

	 	 
추상적인 예측	The photo of Mongolia's present state The photo of Mongolia's state in the future	The photo of Mongolia's present state The photo of Mongolia's state in the future, estimated by an employee Mongolia is facing a harsh reality: its landscapes are turning into desert, its ecosystems are crumbling, and its communities are being uprooted due to climate change. In response, Greentree is launching a vital tree-planting initiative across the nation. This right image depicts the future state of Mongolia, as predicted by AI(employee). According to the AI(employee), planting numerous trees could transform this desertification into grasslands. These trees will help purify the air, potentially lowering Mongolia's temperature. This change will create the conditions to allow many species of animals and plants to thrive. Imagine how your help and donations can change the future of Mongolia. We are not just planting trees; we're also bringing life back to the land and helping to heal our planet. Your donation can make a real difference. Help us transform Mongolia and restore hope to its people.

3) 실험 3A 연구 결과

신뢰도와 기부 의도에 대한 광고 제작자 유형(AI vs. 인간)과 광고 유형(구체적인 예측 vs. 추상적인 예측 vs. 비예측) 각각의 주효과와 둘 간의 상호작용 효과를 검증하기 위해 이원 배치 분산분석을 실시하였다. 광고 제작자 유형의 주효과는 신뢰도와 기부의도에 유의미한 영향을 미치지 않는 것으로 나타났다(신뢰도: $F(1,196) = .001, p = .974$, 기부 의도: $F(1,196) = .492, p = .484$). 광고 유형의 주효과는 신뢰도에 유의미한 영향을 미치지 않았으나($F(1,196) = 1.401, p = .249$), 기부 의도에는 유의미한 영향을 미치는 것으로 나타났다($F(1,196) = 3.691, p = .027$).

광고 제작자 유형과 광고 유형의 상호작용 효과를 살펴본 결과, 신뢰도와 기부 의도에 대한 상호작용 효과는 유의미한 것으로 나타났다(신뢰도: $F(2,196) = 4.684, p = .01$, 기부 의도: $F(2,196) = 5.724, p = .004$)(<그림 5> 참조). 상호작용 효과가 어떻게 나타났는지 본페로니 다중 비

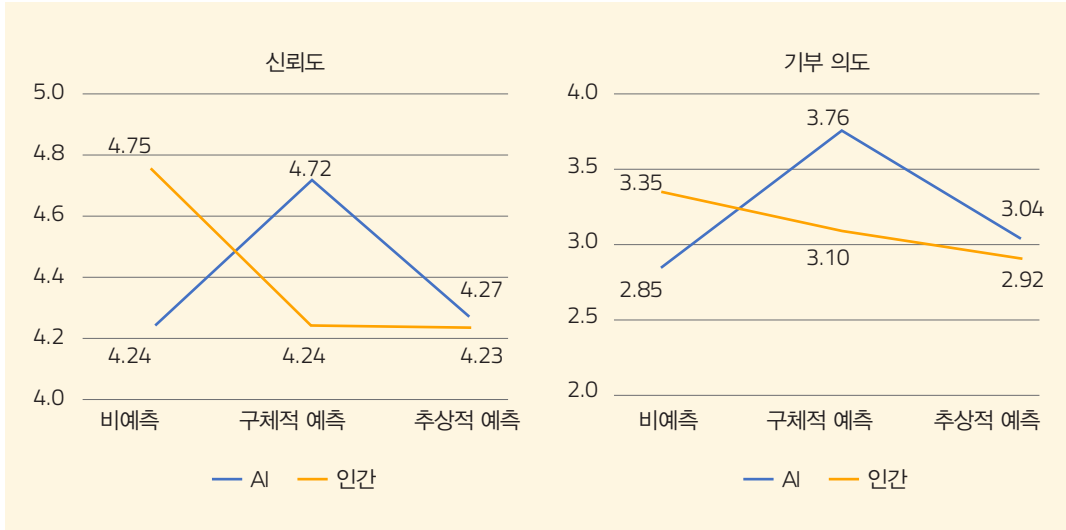
교(Bonferroni's multiple comparison)를 통해 확인하였다(〈표 4〉 참조). 비예측 광고에서는 신뢰도에 대해 인간($M=4.75$, $SD=.162$)과 AI($M=4.24$, $SD=.162$) 간의 유의미한 차이가 났으며, 인간이 만든 광고가 AI가 만든 광고보다 더 신뢰도가 높다고 인식하는 것으로 나타났다($p=.026$). 구체적인 예측을 제시하는 경우에는 AI와 인간 간의 유의미한 차이가 나타났으며($p=.039$), AI가 제작한 광고에 대한 신뢰도가 인간이 제작한 광고보다 더 높은 것으로 나타났다($M_{AI}=4.72$, $SD=.167$, $M_{인간}=4.24$, $SD=.162$). 반면에 추상적인 예측을 제시하는 광고의 경우에는 인간과 AI간의 유의미한 차이가 없는 것으로 나타났다($p=.846$).

기부 의도에 대해서도 비예측 광고인 경우에는 인간이 제작한 광고의 경우에 AI가 제작한 경우보다 더 높은 것으로 나타났다($M_{인간}=3.35$, $SD=.169$, $M_{AI}=2.86$, $SD=.169$, $p=.041$). 구체적인 예측을 제시한 광고의 경우에는 AI가 제작한 광고($M=3.76$, $SD=.174$)에 대한 기부 의도가 인간이 제작한 광고($M=3.1$, $SD=.169$)보다 더 높은 것으로 나타났다($p=.007$). 반면에 추상적인 예측을 제시하는 광고의 경우에는 AI와 인간 간의 유의미한 차이가 없는 것으로 나타났다($p=.598$).

즉, 실험 1, 2와 마찬가지로 비예측의 경우에는 알고리즘 기피 현상이 일어났기 때문에 H1은 지지되었다. 또한, AI가 잘할 것으로 기대하는 예측을 구체적으로 제시하는 경우에는 인간 제작자와 AI 제작자 간의 유의미한 차이가 나타났으며, 구체적으로 알고리즘 기피 현상이 사라지고 인간이 예측한 광고보다 AI의 경우에 더 높은 신뢰도와 기부 의도를 보였다. 반면에 추상적으로 예측을 제시하는 경우에는 인간과 AI 간의 신뢰도와 기부 의도 차이가 유의미하지 않는 것으로 나타나 H4와 H5는 지지되었다.

이전 실험과 마찬가지로 신뢰도의 매개 효과를 검증하기 위해 Hayes(2017)가 제안한 부트스트랩 방법(Process model 8)을 실시하였다. 광고 제작자 유형(X)을 독립 변인, 광고 유형(W)을 조절 변인으로 설정하여 검증한 결과, 매개 변인인 신뢰감(M)이 추가된 모형에서 광고 제작자 유형과 광고 유형의 상호작용이 기부 의도에 미치는 간접 효과는 95%의 신뢰구간 내에서 유의하게 나타나($Effect=.2418$, $S.E.=.0840$, $LLCI=.0893$, $ULCI=.4165$) 신뢰감의 매개 효과를 확인하였다. 이에 따라 H3도 지지되었다.

〈그림 5〉



〈표 4〉

	광고 유형	광고 제작자 유형	표본 수	평균	표준 오차	유의 확률
신뢰도	비예측	AI	34	4.24	.162	.026
		인간	34	4.75	.162	
	예측(concrete)	AI	32	4.72	.167	.039
		인간	34	4.24	.162	
	예측(abstract)	AI	35	4.27	.160	.846
		인간	33	4.23	.165	
기부 의도	비예측	AI	34	2.86	.169	.041
		인간	34	3.35	.169	
	예측(concrete)	AI	32	3.76	.174	.007
		인간	34	3.10	.169	
	예측(abstract)	AI	35	3.04	.167	.598
		인간	33	2.92	.172	

4) 실험 3B 설계 및 절차

실험 3B는 실험3A와 동일하게 2(AI vs. 인간)×3(구체적인 예측 vs. 추상적인 예측 vs. 비예측)으로 집단 간 설계로 구성하였다. 실험 3B에서는 인간의 미래를 예측하였고, 실험 자극물은 〈그림 5〉이다. 실험 3B도 Prolific을 통해서 참가자들을 모집하였으며, 여섯 가지의 조건 중 무작위로 한 조건에 할당되었다.

실험 3B에 사용될 자극물의 적절성을 판단하기 위해 AI의 경우에는 ‘위 광고는 누가 제작하였습니까?’라는 질문을 통해서 조작 점검을 하였다. 110명 중 2명을 제외하고는 참가자들이 제작자를 정확히 응답하여, 광고 유형의 조작이 성공적으로 이루어졌음을 확인할 수 있었다. 측정 문항은 이전 실험과 동일하게 신뢰도($\alpha = .703$)와 기부 의도($\alpha = .779$)를 측정하였다. 총 240명의 참가자들이 본 실험에 참가하였으나 성실하게 답변하지 않았거나 조작 점검에 실패한 참가자들을 제외하고 총 216명(남성=50.2%, M_{age} =38.4세, 백인=71.2%)를 대상으로 분석하였다.

〈그림 5〉

	AI	인간
비예측	 The photo of Haitian children, created by AI	 A The photo of Haitian children
	Right now, Haitian children are struggling with severe acute malnutrition, exacerbated by critical water shortages and inadequate sanitation. Your support and donations are vital. By coming together, we can pull these children out of malnutrition's harsh grip. Please act now—your contribution can truly save lives. With your support and contributions, we can save children from acute malnutrition.	

구체적인 예측	The photo of Haitian children taken in 2024 The photo of Haitian children in 2029, estimated by AI The photo of Haitian children taken in 2024 The photo of Haitian children in 2029, estimated by an employee Right now, Haitian children are struggling with severe acute malnutrition, exacerbated by critical water shortages and inadequate sanitation. Your support and donations are vital. By coming together, we can pull these children out of malnutrition's harsh grip. This picture depicts how AI(employee) predicts children will look in five years. According to the AI's prediction(employee's prediction), children will have received approximately 40 essential nutrients. the AI(employee) predicts that these children will be able to boost their immune systems by 56% and be less vulnerable to about 135 viruses in 5 years with your help. Imagine the difference your donation could make for these children. With your support and contributions, we can save children from acute malnutrition.
추상적인 예측	The photo of Haitian children The photo of Haitian children in the future, estimated by AI The photo of Haitian children The photo of Haitian children in the future, estimated by an employee Right now, Haitian children are struggling with severe acute malnutrition, exacerbated by critical water shortages and inadequate sanitation. Your support and donations are vital. By coming together, we can pull these children out of malnutrition's harsh grip. This picture depicts how AI(employee) predicts children will look in the future. According to the AI(employee)'s prediction, children will have grown up to be healthy. AI(employee) predicts that these children will be able to boost their immune systems and be less vulnerable to viruses. Imagine the difference your donation could make for the children. With your support and contributions, we can save children from acute malnutrition.

5) 실험 3B 연구 결과

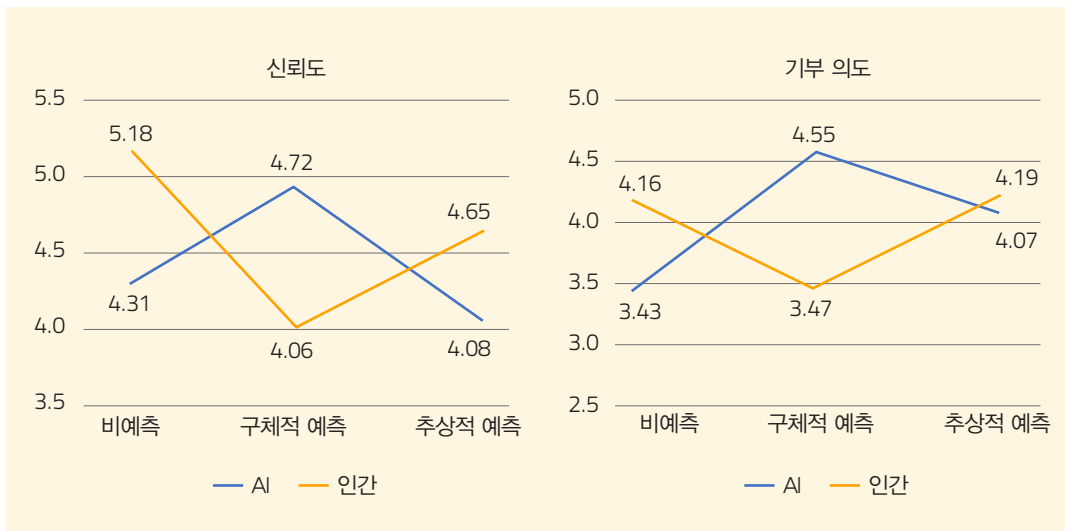
신뢰도와 기부 의도에 대한 광고 제작자 유형(AI vs. 인간)과 광고 유형(구체적인 예측 vs. 추상적인 예측 vs. 비예측) 각각의 주효과와 둘 간의 상호작용 효과를 검증하기 위해 분산분석을 실시하였다. 광고 제작자의 주효과는 신뢰도($F(1,210)=2.193, p=.140$)와 기부 의도($F(1,210)=.107, p=.744$)에는 유의미하지 않는 것으로 나타났다. 광고 유형의 주효과는 신뢰도와 기부 의도 모두 다 유의미한 영향을 미치지 않는 것으로 나타났다(신뢰도: $F(2,210)=2.023, p=.135$, 기부 의도: $F(2,210)=.715, p=.490$).

광고 제작자 유형과 광고 유형에 대한 상호작용 효과를 살펴보면, 신뢰도와 기부 의도 모두 상호작용 효과가 유의미한 것으로 나타났다(신뢰도: $F(2,210)=6.956, p=.001$, 기부 의도: $F(2,210)=5.234, p=.006$)(<그림 6> 참조). 상호작용 효과가 어떻게 나타났는지 본페로니 다중 비교(Bonferroni's multiple comparison)를 통해 확인하였다(<표 5> 참조). 비예측 광고에 대해서는 인간($M=5.18, SD=.210$)이 제작한 광고에 대해 AI($M=4.31, SD=.219$)가 제작한 광고보다 신뢰도가 더 높은 것으로 나타났다($p=.005$). 추상적인 예측을 보여주는 광고에서는 AI와 인간 간의 신뢰도에 대한 유의미한 차이가 없는 것으로 나타났다($p=.067$). 그러나 실험 참가자들은 구체적 예측 조건에 서는, AI가 제작한 광고($M=4.71, SD=.21$)를 인간이 제작한 광고($M=4.06, SD=.222$)보다 더 신뢰하였다($p=.034$).

비예측 광고에 대해서는 광고 제작자 유형에 따라 신뢰도 차이가 유의미하게 나타났고($M_{AI}=3.43, SD=.289, M_{인간}=4.16, SD=.278$), 기부 의도는 유의수준 10% 이내에서 차이를 보였다($p=.071$). 추상적인 예측을 보여주는 광고에 대해서는 AI와 인간 간의 유의미한 차이가 없는 것으로 나타났다($p=.761$). 반면에 구체적인 예측을 보여주는 경우에는 AI($M=4.55, SD=.278$)가 제작한 광고에 대한 기부 의도가 인간($M=3.47, SD=.294$)보다 더 높은 것으로 나타났다($p=.008$). 이에 따라, 구체적으로 예측하는 경우에는 인간 제작자와 AI 제작자 간의 유의미한 차이가 나타났으며, AI가 제작한 경우에 신뢰도와 기부 의도가 인간이 제작한 경우보다 더 높은 것으로 나타났다. 반면에 추상적으로 예측을 제시하는 경우에는 광고 제작자 유형 간의 유의미한 차이가 없는 것으로 나타나 H4와 H5가 지지되었다.

다음으로, 신뢰감의 매개 효과를 검증하기 위해 Hayes(2017)가 제안한 부트스트랩 방법(Process model 8)을 실시하였다. 광고 제작자 유형(X)을 독립 변인, 광고 유형(W)을 조절 변인으로 설정하여 검증한 결과, 매개 변인인 신뢰감(M)이 광고 제작자 유형과 광고 유형의 상호작용이 기부 의도에 미치는 간접 효과는 95%의 신뢰구간 내에서 유의하게 나타나($Effect = .6220, S.E. = .1651, LLCI = .3057, ULCI = .9576$) 신뢰감의 매개 효과를 확인하였다. 이에 따라 H3도 지지되었다.

〈그림 6〉



〈표 5〉

	광고 유형	광고 제작자 유형	표본 수	평균	표준 오차	유의 확률
신뢰도	비예측	AI	35	4.31	.219	.005
		인간	38	5.18	.210	
	예측(concrete)	AI	38	4.71	.210	.034
		인간	34	4.06	.222	
	예측(abstract)	AI	37	4.08	.213	.067
		인간	34	4.65	.222	
기부 의도	비예측	AI	35	3.43	.289	.071
		인간	38	4.16	.278	
	예측(concrete)	AI	38	4.55	.278	.008
		인간	34	3.47	.294	
	예측(abstract)	AI	37	4.07	.281	.761
		인간	34	4.19	.294	

6) 실험 3 논의

실험 3은 실험 1과 실험 2의 결과를 재현하고 확장하기 위해서 진행되었으며, 주요 연구 결과는 다음과 같다. 첫째, 실험 1과 실험 2와 동일하게 비예측 광고의 경우에는 AI보다 인간이 만든 광고에 대해 신뢰감과 기부 의도가 더 높았다. 반면에 구체적인 예측을 보여주는 광고는 AI가 제작했다고 알려주는 경우, 인간이 제작했다고 알려주는 경우보다 소비자들은 더 높은 신뢰도와 기부 의도를 보였다. 하지만 예측이 추상적인 경우에는, AI 제작자가 만든 광고와 인간 제작자가 만든 광고 사이에 유의미한 차이가 나타나지 않았다. 이는 AI를 활용한 예측 광고를 제작할 경우, 구체적이고 객관적인 내용을 전달할 때 가장 효과적인 것임을 시사한다. 반면에 인간이 제작한 예측 광고는 구체적인 내용으로 기술될지라도 신뢰도가 크게 증가하지 않았다.

6. 요약 및 시사점

본 논문은 AI가 제작한 광고에 대한 소비자 태도를 기부 영역에서 살펴본 연구이다. 흔히 창작의 특성을 지닌 업무에 대해 소비자들은 알고리즘 기피 현상을 보이는 것으로 보고되어 왔는데, 이것이 언제나 그러한지, 알고리즘 수용 현상, 혹은 알고리즘 존중 현상으로 바뀌는 조건이 있을지 탐색하였다. 실험 1과 실험 2, 실험 3A와 실험 3B를 통해 발견한 내용은 아래와 같다.

첫째, 비예측 광고는 AI보다 인간이 제작한 기부 광고에 대해 소비자들은 더 높은 신뢰도와 기부 의도를 보였다. 이 연구 결과는 기부 영역이 아닌 일반적인 광고 분야에서 보고되어 왔던 선행 연구의 방향성을 지지한다.

둘째, 미래를 예측하는 광고에서 소비자들의 알고리즘 기피 현상은 감소했다. 선행 연구는 소비자들이 AI가 수행한 업무가 객관적이라고 인식할 때, 그에 대한 선호도가 증가함을 발표했다(Wu & Wen, 2021). 본 연구는 비록 광고라 할지라도 예측이 포함된 광고는 소비자들이 객관적인 업무로 인식하고 따라서 AI가 제작한 예측 광고에 대한 신뢰도와 기부 의도가 인간 제작자가 만든 예측 광고보다 높게 나타난 것으로 보인다. AI가 인간보다 객관적인 업무를 더 정확하고 유능하게 처리할 것이라고 믿는 AI에 대한 인식이 AI가 제작한 예측 광고에 대한 신뢰도와 기부 의도에 영향을 미친 것으로 해석할 수 있다.

셋째, 이미지를 활용한 예측 메시지의 경우에도 알고리즘 기피 현상이 약화되는 것으로 나타났다.

넷째, 소비자들의 알고리즘 수용 현상은 추상적인 예측보다 구체적인 예측의 경우에 강하게 드러났다. 즉, AI가 구체적인 예측 광고를 제시한 경우에 인간이 제시한 경우보다 소비자들은 더 높은 신뢰도와 기부 의도를 보였다.

본 연구의 이론적, 실무적 공헌점은 다음과 같다.

첫째, 기존 연구들은 AI에 대한 소비자들의 부정적 태도, 알고리즘 기피 현상이 발생하는 상황에 초점을 맞추어 연구를 진행했으나 본 연구는 어떻게 소비자들의 알고리즘 기피 현상이 완화되는지를 살펴보았다. 이를 통해 소비자들의 알고리즘 기피와 선호에 대한 이해를 확장하고 있다. 이는 앞으로 가속하여 진행될 생성형 AI의 도입과 활용에 대한 소비자들의 태도를 이해하는 데 도움을 주고 비영

리 조직을 포함한 기업들이 어떻게 대응해야 할지 전략적 방안을 모색하는 데 시사점을 제공한다.

둘째, 예측을 보여주는 광고에 대한 소비자 인식을 살펴보았다. 미래 예측은 장기적으로 우리의 삶을 변화시킬 것이라는 기대감을 주기에 현재의 행동에도 영향을 미친다(Ye et al., 2015). 이런 중요성에도 불구하고 예측을 포함하는 메시지에 대한 소비자, 시민, 기부자들에 대한 소비 행동 관점의 연구는 제한적이다. 본 연구는 예측을 보여주는 광고에 대한 효과성을 탐구하였다는 점에서 의의가 있다.

셋째, 본 연구는 고정 관념 이론과 머신 휴리스틱 이론을 바탕으로 AI에 대한 소비자 인식이 기부 행동에 미치는 영향을 분석하였다. 선행 연구들이 기업의 추천 서비스나 상업 광고를 분석한 것과 달리, 본 연구는 사람들의 AI, 머신에 대한 판단과 인식, 태도에 대한 논의를 기부 행동으로 연결하여 탐구했다는 점에서 해당 이론의 적용 범위를 확장하였다.

본 연구를 통해 다음과 같은 실무적 시사점을 도출할 수 있다.

첫째, 비예측 형태의 기부 광고에서는 AI보다 인간이 제작한 광고가 더 높은 신뢰도와 기부 의도를 유발하였다. 따라서 일반적인 사회적 메시지를 전달할 때에는 인간적인 감성과 휴머니즘에 기반한 접근이 더 효과적일 수 있음을 시사한다.

둘째, 예측 광고의 경우 AI가 제작했을 때 알고리즘 기피 현상이 감소했으며, 인간이 제작한 광고보다 높은 신뢰도와 기부 의도를 보였다. 따라서 예측과 분석에 기반한 캠페인에서는 AI의 객관성을 부각시키는 것이 중요하다. 예를 들어 데이터 기반의 예측을 통해 미래 상황을 설명하는 기부 광고에서는 AI를 활용하는 것이 기부자(소비자)를 설득하는 데 도움이 될 수 있다.

셋째, 텍스트가 아닌 이미지를 활용한 예측 광고에서도 AI의 알고리즘 기피 현상이 완화되었다. 이는 일반적으로 소비자들이 인식하기에 창의성이 요구되는 시각적 요소에서도 AI의 예측 기능이 더 높은 신뢰도를 이끌어 낼 수 있음을 시사한다. 시각적 이미지가 포함된 예측 광고에서도 AI를 활용하는 것이 효과적인 광고 메시지 전달에 기여할 수 있음을 보여준다.

넷째, 예측 광고가 추상적인 형태보다 구체적인 형태로 제시될 때 소비자들은 알고리즘 선호 현상을 보였다. 따라서 AI를 활용한 예측 기부 광고에서는 구체적인 데이터와 명확한 수치를 제시하면 더 효과적일 것이다. 기존의 기부 광고에서 따뜻한 메시지라는 감성적 특성을 주로 활용해 왔지만, AI, 미래, 예측, 분석, 객관적, 데이터와 같은 조합으로 메시지를 제시한다면 감성적 특성과 다른 새로운

커뮤니케이션 전략이 필요할 수 있음을 시사한다. 본 연구 결과와 기부 단체의 브랜드 개성, 기부자들의 개인적 특성, 따뜻함과 공감의 결합됐을 때 어떤 효과를 야기할지 추가적인 연구가 필요하다.

마지막으로, 본 연구는 기부금이 특정 문제에 적용되어 일으킬 미래 변화 결과를 구체적으로 보여주는 경우에는 AI의 활용이 오히려 더 효과적일 수 있음을 시사한다. 이는 소비자들은 AI가 제공하는 예측을 데이터 기반의 객관적 정보로 인식하기 때문에, 더 깊은 신뢰감을 줄 수 있다. 예를 들어 기부금이 사용될 구체적인 프로젝트나 그로 인해 달성될 미래 변화를 AI가 예측해 시각화한다면, 소비자들은 해당 예측을 더 현실적이고 신뢰할 수 있는 정보로 받아들일 가능성이 높다. 이는 AI 기술이 단순히 광고 제작에 사용되는 것 이상의 가치를 가지며, 소비자에게 구체적이고 설득력 있는 미래 전망을 제공함으로써 기부 의도를 강화할 수 있음을 시사한다.

본 연구는 앞에서 언급한 바와 같은 학문적, 실무적 의미를 내포한다. 그러나 다음과 같은 한계점 및 향후 연구에 대한 필요성도 있다.

첫째, 광고 제작자 유형과 소비자 태도가 조절될 수 있는 변수들이 있다. 예를 들면, 선행 연구는 수혜자들의 사진이 행복한 모습인지, 슬픈 모습인지에 따라 기부자의 태도가 달라질 수 있으며, 수혜자가 몇 명인지에 따라서도 달라질 수 있음을 확인하였다(Li & Yin, 2022). 따라서, 수혜자의 미래 중 긍정, 부정 요소 중에 어떤 모습을 보여주느냐에 따라서 연구 결과에 영향을 미치는지 후속 연구가 필요하다.

또한, AI와 관련된 많은 선행 연구는 AI의 의인화에 대해 다루고 있다. 이에 따라 예측 광고를 창작한 AI의 의인화 정도가 소비자(기부자)들의 태도에 어떤 영향을 미칠지에 대한 연구도 필요하다.

참고 문헌

- 임지선 (2023, 07.04). 주현영이 AI로 광고를 만들면?…LGU+의 색다른 시도, *한겨레*.
https://www.hani.co.kr/arti/economy/economy_general/1098624.html
- Ahn, J., Kim, J., & Sung, Y. (2022). The effect of gender stereotypes on artificial intelligence recommendations, *Journal of Business Research*, 141, 50–59.
- Arango, L., Singaraju, S. P., & Niininen, O. (2023). Consumer responses to AI-generated charitable giving ads, *Journal of Advertising*, 52(4), 486–503.
- Brüns, J. D., & Meißner, M. (2024). Do you create your content yourself? Using generative artificial intelligence for social media content creation diminishes perceived brand authenticity, *Journal of Retailing and Consumer Services*, 79, 103790.
- Bonezzi, A., & Ostinelli, M. (2021). Can algorithms legitimize discrimination?. *Journal of Experimental Psychology: Applied*, 27(2), 447.
- Campbell, C., Plangger, K., Sands, S., Kietzmann, J., & Bates, K. (2022). How deepfakes and artificial intelligence could reshape the advertising industry: The coming reality of AI fakes and their potential impact on consumer behavior, *Journal of Advertising Research*, 62(3), 241–251.
- Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion, *Journal of Marketing Research*, 56(5), 809–82.
- Chaiken, S., & Maheswaran, D. (1994). Heuristic processing can bias systematic processing: effects of source credibility, argument ambiguity, and task importance on attitude judgment, *Journal of personality and social psychology*, 66(3), 460.
- Chiang, T. H., Liao, C. S., & Wang, W. C. (2022). Impact of artificial intelligence news source credibility identification system on effectiveness of media literacy education, *Sustainability*, 14(8), 4830.
- De Groot, J. I. M. (2022). The personalization paradox in Facebook advertising: The mediating effect of relevance on the personalization – brand attitude relationship and the moderating effect of intrusiveness, *Journal of Interactive Advertising*, 22(1), 57–74.
- Gao, Y., & Liu, H. (2023). Artificial intelligence-enabled personalization in interactive marketing: a customer journey perspective, *Journal of Research in Interactive Marketing*, 17(5), 663–680.

- Gotlieb, J. B., & Sarel, D. (1991). Comparative advertising effectiveness: The role of involvement and source credibility. *Journal of advertising*, 20(1), 38–45.
- Hautz, J., Füller, J., Hutter, K., & Thürridl, C. (2014). Let users generate your video ads? The impact of video source and quality on consumers' perceptions and intended behaviors. *Journal of Interactive Marketing*, 28(1), 1–15.
- Hofeditz, L., Mirbabaie, M., Holstein, J., & Stieglitz, S. (2021, June). Do You Trust an AI-journalist? A Credibility Analysis of News Content with AI-Authorship. In *ECIS*.
- Holthöwer, J., & van Doorn, J. (2023). Robots do not judge: service robots can alleviate embarrassment in service encounters. *Journal of the Academy of Marketing Science*, 51(4), 767–784.
- Hovland, C. I., Janis, I. L., & Kelley, H. H. (1953). *Communication and persuasion; Psychological studies of opinion change*. New Haven, CT: Yale University Press.
- Jung, M., & Seiter, M. (2021). Towards a better understanding on mitigating algorithm aversion in forecasting: An experimental study. *Journal of Management Control*, 32(4), 495–516.
- Im, H., & Lee, G. (2023). Do consumers always believe humans create better boxes than AI? The context-dependent role of recommender creativity. *International Journal of Retail & Distribution Management*, 51(8), 1045–1060.
- Ismagilova, E., Slade, E., Rana, N. P., & Dwivedi, Y. K. (2020). The effect of characteristics of source credibility on consumer behaviour: A meta-analysis. *Journal of Retailing and Consumer Services*, 53, 101736.
- Kelly, Harold H. (1972). Attributions in Social Interaction, in *Attribution: Perceiving the Causes of Behavior*, E.E. Jones, D.E. Kanouse, H.H. Kelly, R.E.
- Khachaturian, J. and Morganosky, M.A. (1990), Quality perceptions by country of origin, *International Journal of Retail and Distribution Management*, 18(5), 21–30.
- Kim, J., Giroux, M., & Lee, J. C. (2021). When do you trust AI? The effect of number presentation detail on consumer trust and acceptance of AI recommendations. *Psychology & Marketing*, 38(7), 1140–1155.
- Kolbl, Z., Diamantopoulos, A., Arslanagic-Kalajdzic, M. and Zabkar, V. (2020), Do brand warmth and brand competence add value to consumers? A stereotyping perspective. *Journal of Business*

- Research*, 118, 346–362.
- Lee, J. K., Lee, S. Y., & Hansen, S. S. (2017). Source credibility in consumer-generated advertising in YouTube: The moderating role of personality. *Current Psychology*, 36, 849–860.
- Li, M. R., & Yin, C. Y. (2022). Facial expressions of beneficiaries and donation intentions of potential donors: Effects of the number of beneficiaries in charity advertising. *Journal of Retailing and Consumer Services*, 66, 102915.
- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103.
- Lou, C., & Yuan, S. (2019). Influencer marketing: How message value and credibility affect consumer trust of branded content on social media. *Journal of interactive advertising*, 19(1), 58–73.
- Lv, L., & Huang, M. (2024). Can personalized recommendations in charity advertising boost donation? The role of perceived autonomy. *Journal of Advertising*, 53(1), 36–53.
- Muda, M., & Hamzah, M. I. (2021). Should I suggest this YouTube clip? The impact of UGC source credibility on eWOM and purchase intention. *Journal of Research in Interactive Marketing*, 15(3), 441–459.
- Nass, C., Moon, Y. and Green, C. (1997). Are machines gender-neutral? Gender-stereotypic responses to computers with voices. *Journal of Applied Social Psychology*, 27, 864–876.
- Qin, X., & Jiang, Z. (2019). The impact of AI on the advertising process: The Chinese experience. *Journal of Advertising*, 48(4), 338–346.
- Shah, N., Engineer, S., Bhagat, N., Chauhan, H., & Shah, M. (2020). Research trends on the usage of machine learning and artificial intelligence in advertising. *Augmented Human Research*, 5, 1–15.
- Sundar, S. S., & Kim, J. (2019, May). Machine heuristic: When we trust computers more than humans with our personal information. In *Proceedings of the 2019 CHI Conference on human factors in computing systems* (pp. 1–9).
- Whyte, C. (2020). Deepfake news: AI-enabled disinformation as a multi-level public policy challenge. *Journal of cyber policy*, 5(2), 199–217.
- Wu, L., Dodoo, N. A., & Wen, T. J. (2024). Disclosing AI's Involvement in Advertising to Consumers: A Task-Dependent Perspective. *Journal of Advertising*, 1–19.
- Wu, L., & Wen, T. J. (2021). Understanding AI advertising from the consumer perspective: What factors determine consumer appreciation of AI-created advertisements?. *Journal of Advertising Research*, 61(2), 133–146.

- Yang, J., Chunterawong, P., Lee, H., Tian, Y., & Chock, T. M. (2023). Human versus virtual influencer: the effect of humanness and interactivity on persuasive CSR messaging. *Journal of Interactive Advertising*, 23(3), 275–292.
- Ye, N., Teng, L., Yu, Y., & Wang, Y. (2015). “What’ s in it for me?” : The effect of donation outcomes on donation behavior. *Journal of Business Research*, 68(3), 480–486.
- Zhou, Y., Fei, Z., He, Y., & Yang, Z. (2022). How human – chatbot interaction impairs charitable giving: the role of moral judgment. *Journal of Business Ethics*, 178(3), 849–865.