

스마트 음성 비서의 유머 스타일이 지각된 도덕성에 미치는 영향

The impact of smart voice assistants' humor styles on perceived morality

차무화¹⁾ • 박현정^{2)*}

충북대학교 국제경영학과 박사과정¹⁾ • 충북대학교 국제경영학과 교수^{2)*}

Che, Mao Hua¹⁾ • Park, Hyun Jung^{2)*}

Department of International Business, Chungbuk National University^{1),2)}

Abstract

Smart voice assistants have increasingly integrated humor into their interactions with consumers. However, as these technologies become more human-like, concerns regarding the moral quality of their humor emerge, given that inappropriate humor can demean others and lead to perceptions of low morality.

This study defines consumers' moral judgment of smart voice assistants' behavior as perceived morality and investigates how different humor styles influence this perception. We conducted a survey with users, categorizing smart voice assistants into five distinct humor styles and randomly assigning one style to each participant. The results indicated that affiliative humor, self-enhancing humor, and self-defeating humor positively influenced consumer attitudes and perceived morality, while aggressive humor yielded a negative effect. Additionally, we found that the influence of humor styles on consumer attitudes can be mediated by perceived morality. Implicit theory significantly moderates the effects of affiliative and aggressive humor on perceived morality and attitudes. Incremental theorists demonstrated that stronger affiliative humor correlated with higher perceived morality and more favorable attitudes, whereas entity theorists noted that stronger aggressive humor was associated with lower perceived morality and less favorable attitudes.

This study provides a theoretical foundation for understanding the effective use of humor in smart voice assistants and offers practical guidelines for developers and marketers regarding the careful and appropriate application of humor based on consumer characteristics.

Keywords: Consumer attitude, Humor styles, Implicit theory, Perceived morality, Smart voice assistant

I. 서론

4차 산업혁명의 급속한 발전으로 인해 인공지능 기술을 활용한 다양한 서비스 관련 산업이 등장하면서, 상품 마케팅과 판매 등의 서비스가 보편화되고 있다(모은가, 류미현, 2023). 인공지능 기술을 활용한 스마트 음성 비서는 소비자와의 상호작용을 통해 맞춤형(Behera et al., 2024) 및 참

여형(Chung et al., 2020) 서비스를 제공할 수 있다. 그러나 여전히 많은 스마트 음성 비서가 소비자의 기대에 부응하지 못하고 있어(Shreehan et al., 2020), 소비자 태도를 향상하기 위해 스마트 음성 비서 서비스를 어떻게 디자인해야 할지에 대한 고민이 필요한 상황이다.

이전 연구들은 언어 스타일, 이름, 성별, 에이전트의 개인 데이터 이미지, 의도적 지연(Shreehan et al., 2020;

본 논문은 2024 한국생활과학회 하계학술대회에서 발표되었음

* Corresponding author: Park, Hyun Jung

Tel: +82-43-261-2331, Fax: +82-43-273-1451

E-mail: phj@cbnu.ac.kr

© 2024, Korean Association of Human Ecology. All rights reserved.

Toader et al., 2020) 등 다양한 디자인 요소를 조사했다. 그러나 스마트 음성 비서가 구사하는 유머 스타일이 소비자와 관련된 서비스 결과에 미치는 영향에 대해서는 알려진 바가 거의 없다(Thomaz et al., 2020). Schanke et al.(2021)은 유머가 스마트 음성 비서와 상호작용하는 소비자의 태도를 개선할 수 있는지는 불분명하다고 지적하였다. 특히 공격적이거나 비도덕적인 행동은 소비자 태도에 큰 부정적인 영향을 미칠 수 있으므로, 스마트 음성 비서의 부정적인 유머 스타일이 소비자의 태도에 미치는 영향도 고려해야 한다. Shin et al.(2023)은 챗봇이 유머를 사용하여 서비스의 매력과 오락성을 향상시키는지 살펴본 연구에서 친화적 유머와 공격적 유머의 긍정적 역할만 논의했을 뿐, 자기고양과 자기패배적 유머는 아직 아무도 언급하지 않았다.

이에 본 연구에서는 스마트 음성 비서가 사용하는 다양한 유머 스타일(친화적 유머, 자기고양적 유머, 자기패배적 유머, 공격적 유머)이 제품 태도에 미치는 영향을 조사하였다. 특히, 본 연구는 스마트 음성 비서의 유머 스타일이 소비자의 지각된 도덕성을 통해 제품 태도에 어떤 영향을 미치는지를 탐구한다. 스마트 음성 비서는 인간과 의사소통하고 상호작용하도록 인간과 유사하게 설계되었기 때문에, 일반적으로 소비자는 스마트 음성 비서가 사회의 도덕적 규범을 따를 것이라고 기대한다.

선행연구에 따르면, 소비자의 암묵적 이론(implicit theories)은 인공지능 로봇에 대한 인식과 평가에 영향을 미친다(Dang & Liu, 2022). 구체적으로, 증진 이론가들은 로봇에 대한 부정적인 감정 수준을 낮추고, 로봇을 적보다는 파트너로 더 많이 보고, 로봇에 관한 연구를 더 많이 지원하며, 로봇과 상호작용하는 것을 더 선호한다. 또한 로봇이 높은 수준의 마음을 가진다고 생각될 때 로봇의 긍정적인 반응에 대한 증진 이론가의 영향이 더 두드러진다. Kwon과 Nayakankuppam(2015)의 연구에 따르면, 불변 이론가는 큰 노력을 들이지 않고도 빠르게 태도를 형성하고 증진 이론가보다 이러한 태도를 더 오래 유지하는 경향이 있다. 반면, 증진 이론가는 사람의 변화에 대해 더 관대한 견해를 가지고 있다(Hong et al., 1999).

본 연구는 불변 이론가의 고정된 태도와 증진 이론가의 관대한 태도가 스마트 음성 비서가 사용하는 유머 스타일과 관련된 맥락에서도 영향력을 행사할 것인지, 또 사용된 유머 스타일에 따라 어떤 다른 효과가 있는지를 탐구함으로써 암묵적 이론, 의인화 및 유머에 관한 소비자 연구에 기여하고자 한다.

II. 이론적 배경 및 가설 수립

1. 유머 사용의 효과

유머는 웃음, 오락, 또는 재미를 유발하기 위해 의도된 자극으로 정의된다(Warren et al., 2018). 유머는 사회적 관계를 윤활화(lubricate)하며(Duncan & Feisal, 1989), 긴장 상황에서 스트레스(Scott, 2007)와 갈등(Norrick & Spitz, 2008)을 줄이는 데 도움이 된다.

소비자 연구에 따르면 유머는 소비자의 관심을 높이고 광고에 대한 태도(Chung & Zhao, 2003), 브랜드에 대한 태도(Van et al., 2011), 구매 의도(Eisend, 2009) 등에 긍정적인 영향을 미친다. 유머는 경박함을 동반하여 비하 메시지를 부드럽게 하거나 완화시키며, 때로는 도덕적으로 수용하기 쉽게 설명을 변경할 수 있다(Ford et al., 2001). 이는 무례한 광고에 수반되는 윤리적 문제를 어느 정도 덮을 수 있는 요소로 작용할 수 있다(Förster, 2016).

한 연구에서는 여성 비하 유머 광고가 소비자의 지각된 도덕성에 부정적인 영향을 미친다는 결과가 도출되었다(Förster, 2016). 그러나 대부분의 연구에서는 유머가 소비자 행동에 미치는 긍정적 효과를 강조하며, 유머 사용에 따라 제기될 수 있는 도덕적이고 윤리적인 문제는 종종 간과되어 왔다.

일부 연구는 광고에서 유머의 효과에 있어 메시지 특성을 조절 변수로 강조하고 있다. 즉, 유머의 영향력은 광고 메시지의 내용에 따라 달라진다고 주장한다. Beard(2008)는 광고에 묘사된 공격적인 주제와 의도적인 유머에 대해 게시된 300여건의 불만 사항을 조사하였다. 그는 시청자를 불쾌하게 하는 광고에는 종종 의도적인 유머가 포함되어 있지만, 불쾌감과 의도적인 유머의 관계는 크지 않고 유머에 포함된 고유한 공격적인 주제와의 관계가 더 크다는 것을 발견했다. 따라서 부적절한 유머는 윤리적 위반의 대상이 될 수 있으며, 도덕규범을 위반하는 것으로 간주될 수 있다.

스마트 음성 비서가 유머를 활용하는 것은 예상치 못한 방식으로 상호작용을 더욱 흥미롭고 매력적이며 신뢰할 수 있도록 만들어, 비서에게 개성을 부여하는 중요한 요소로 평가된다(Braslavski et al., 2018; Chaves & Gerosa, 2021; Grudin & Jacques, 2019). 유머는 스마트 음성 비서의 인터페이스를 보다 친근하고 관련성 있게 만드는 데 기여하며, 대화형 에이전트 설계에 자주 통합되고 있다(Lopatovska, 2020). 예를 들어, 두 개의 대화형 로봇을 비

교한 실험 연구에서는 외향적인 성격을 드러내고 유머를 사용하도록 프로그래밍된 시스템이 사용자의 인식과 전반적인 즐거움을 향상시키는 것으로 나타났다(Niculescu et al., 2013).

스마트 음성 비서가 네 가지 서로 다른 유머 스타일을 사용할 때, 소비자가 느끼는 윤리적 평가와 서비스에 대한 인식이 달라질 수 있다. 스마트 음성 비서가 실제 사회적 행위자이든 아니든, 의인화로 인해 유머의 적절한 사용에 대한 기대를 포함한 사회적 상호작용의 규범을 따를 것으로 예상된다. Følstad et al.(2018)은 스마트 음성 비서의 유머가 소비자와의 상호작용에서 신뢰를 구축하는 데 기여할 가능성이 있지만, 이는 적절하게 사용될 때만 가능하다고 주장했다.

사람들이 대화에서 유머를 사용하는 이유는 스트레스를 줄이고, 놀라운 상황에 대처하며, 단합과 소속감을 강화하기 위함이다(Lopatovska, 2020). 따라서 유머가 스마트 음성 비서의 기술 설계에 적용되면, 더 매력적이고 개성이 있어야 한다. 본 논문에서는 스마트 음성 비서가 유머를 합리적으로 사용할 경우, 소비자의 도덕적 윤리적 인식에 긍정적인 영향을 미치고, 반대로 유머 사용이 부적절할 경우 부정적인 영향을 미칠 것이라고 가정하였다.

2. 유머 스타일과 소비자 반응

유머 스타일을 이해하는 것은 스마트 음성 비서와 같은 대화형 시스템에서 유머를 효과적으로 활용하기 위한 기초를 제공한다고 볼 수 있다. 유머의 신중한 선택은 유머의 성공적인 적용을 보장하는 데 매우 중요하기 때문이다(Kellaris & Cline, 2007). 선행연구에 따르면 네 가지 유형의 유머 스타일 중 두 가지(친화적 유머와 자기고양적 유머)는 사회 관계 유지에 긍정적인 영향을 미치지만, 나머지 두 가지(자기패배적 유머와 공격적 유머)는 덜 유익하거나 심지어 행복에 해를 끼칠 수 있다(Martin et al., 2003).

친화적 유머(affiliative humor)는 타인을 즐겁게 하도록 사용되나, 해롭지 않은 방식으로 활용되며, 사회적 관계를 형성하고 강화하며 긴장을 줄이는 데 기여한다. 친화적 유머가 스마트 음성 비서에 사용될 때 서비스 만족도를 높일 수 있다는 연구 결과가 있다(Shin et al., 2023). 이 유머의 효과는 소비자들이 스마트 음성 비서를 더 인간적으로 느끼게 하고, 상호작용을 더욱 재미있게 만들어준다. 따라서 이 연구는 스마트 음성 비서와 소비자 간의 상호작용에서 친화적 유머가 음성 비서의 도덕성 인식과 소비자 태도

에 긍정적인 영향을 미치리라 추측한다.

자기고양적 유머(self-enhancing humor)는 인생의 스트레스 상황과 역경에 대해 유머러스한 태도를 유지하는 것을 의미하며, 이는 감정을 조절하고 스트레스를 줄이기 위한 대처 전략으로 기능한다. 자기고양적 유머는 행복, 희망, 낙관적인 태도를 키우는 데 도움을 준다(Thorson et al., 1997). 또한, 자기고양적 유머는 자존감과 양의 상관관계가 있으며, 자존감이 높은 사람은 자기고양적 유머를 사용하는 경향이 있다. 자기고양적 유머에 적응한 사람일수록 주관적 행복감을 느끼는 것으로 나타났다(Yue, 2014). 따라서 본 연구는 스마트 음성 비서가 자기고양적 유머를 사용할 때 사람들이 비서의 도덕성을 더 긍정적으로 인식하고 긍정적인 태도를 형성하는 데 도움이 될 것이라고 예상한다.

반면, 자기패배적 유머(self-defeating humor)는 다른 사람을 즐겁게 하려고 자신을 대상으로 삼는 유머로, 과도하게 자신을 희생하는 경향이 있다. 자기패배적 유머는 공격적 유머와 마찬가지로 부정적인 유머에 속하지만, 스마트 음성 비서의 자기패배적 유머는 서비스 회복 만족도를 높이는 긍정적인 영향을 미친다(Yang et al., 2023). 전통적인 사과 방식보다 자기패배적 유머를 사용할 때 소비자의 감정적 공감을 자극하고, 더 진정한 느낌을 주어 고객의 용서를 받기 쉬워지기 때문이다(Yang et al., 2023). 자기패배적 유머는 따뜻함에 대한 인식을 만들어 소비자 태도에 긍정적인 영향을 미친다(Jiang et al., 2019). 그러므로 스마트 음성 비서가 자기패배적 유머를 사용할 때 도덕성 인식과 제품 태도를 증가시킬 것으로 보인다.

마지막으로, 공격적 유머(aggressive humor)는 다른 사람을 비하하기 위해 조롱과 놀림을 사용하는 유머로, 그 도는 상대방을 해치거나 비난하는 데 있다(Martin et al., 2003). 공격적 유머는 스마트 음성 비서에서 사용될 때 서비스 상호작용에서 친화적 유머보다 덜 효과적인 것으로 나타났다(Shin et al., 2023). 선행연구에 따르면 공격적인 유머는 관계 만족도에 부정적인 영향을 미칠 수 있다(Campbell et al., 2008). 서비스 직원이 공격적인 유머를 사용하면 서비스에 대한 불만족을 초래할 수 있으며(McColl-Kennedy & Sparks, 2003), 광고에서 불쾌하거나 부적절한 유머가 사용되면 소비자들이 해당 기업에 부정적인 태도를 가지게 된다(Warren & McGraw, 2016). 따라서 스마트 음성 비서가 공격적인 유머를 사용할 때 소비자들이 비서의 도덕성을 낮게 평가하고 부정적인 태도를 형성할 것이다. 이에 따라 다음과 같은 가설을 수립하였다.

H1. 스마트 음성 비서의 a) 친화적 유머, b) 자기고양적 유머, c) 자기패배적 유머는 스마트 음성 비서에 대한 소비자의 지각된 도덕성에 정(+)의 영향을 미치고, d) 공격적 유머는 지각된 도덕성에 부(-)의 영향을 미칠 것이다.

H2. 스마트 음성 비서의 a) 친화적 유머, b) 자기고양적 유머, c) 자기패배적 유머는 소비자 태도에 정(+)의 영향을 미치고, d) 공격적 유머는 소비자 태도에 부(-)의 영향을 미칠 것이다.

기업, 브랜드나 서비스 제공자에 대한 지각된 도덕성은 현재 비즈니스 세계에서 매우 중요하다. 브랜드에 대한 지각된 도덕성은 고객 충성도, 신뢰도, 감정적 연결, 제품 품질 등에 대한 인식을 높여 브랜드에 대한 태도를 개선한다(Batool et al., 2023). 연구들은 지각된 도덕성이 브랜드 충성도(Markovic et al., 2018; Singh, 2012), 브랜드 신뢰도(Singh, 2012), 구매 의도(Brunk, 2012), 제품 또는 서비스 품질(Markovic et al., 2018; Sierra et al., 2017), 그리고 긍정적인 입소문(Markovic et al., 2018)을 증가시킬 수 있음을 보여주었다. 이러한 결과는 지각된 도덕성이 브랜드나 기업에 많은 이점을 가져다줄 수 있음을 나타낸다. 따라서 스마트 음성 비서에 대한 소비자의 도덕성 인식은 해당 비서에 대한 태도에 긍정적인 영향을 미칠 것이다. 즉, 스마트 음성 비서의 유머 스타일이 소비자 태도에 미치는 영향은 지각된 도덕성을 통해 매개될 것으로 보인다.

H3. 스마트 음성 비서에 대해 지각된 도덕성은 소비자 태도에 정(+)의 영향을 미칠 것이다.

H4. 스마트 음성 비서의 a) 친화적 유머, b) 자기고양적 유머, c) 자기패배적 유머, d) 공격적 유머가 소비자 태도에 미치는 영향은 지각된 도덕성을 통해 매개될 것이다.

3. 암묵적 이론

암묵적 이론(implicit theories)은 사람들이 세상을 이해하기 위해 사용하는 인간 특성의 본질에 대한 일반적인 믿음을 의미한다(Chiu et al., 1997). 이러한 이론은 인간 본성에 대한 믿음을 두 가지 뚜렷한 범주로 구분한다: 불변 이론(entity theory)과 증진 이론(incremental theory). 불변 이론가들은 개인의 속성과 자질이 안정적이고 고정된 특성이라고 믿지만, 증진 이론에 동의하는 사람들은 이러한

속성과 자질이 상대적으로 가변적이고 역동적이라고 믿는다(Song, 2024).

이모티콘은 유머의 매개체 역할을 하며, 유머러스한 이모티콘을 사용하는 챗봇이 사용하지 않는 경우에 비해 서비스 실패 상황에서도 소비자의 재사용 의도를 높일 수 있다. 또 증진 이론을 가진 사람들은 불변 이론을 가진 사람들보다 이러한 효과를 더 강하게 느끼는데(Liu et al., 2023), 이들은 기술이 발전함에 따라 스마트 음성 비서가 점점 더 인간과 유사한 능력을 개발할 것이라고 믿는다(Croes, 2021). 따라서 스마트 음성 비서가 친화적인 유머를 사용할 때, 증진 이론가들은 불변 이론가들보다 해당 비서를 보다 의인화하고 도덕성이 높은 존재로 인식할 가능성이 크다.

반면, 불변 이론가들은 서비스 제공자가 한 번 부정적인 결과를 초래하면 이를 앞으로도 지속적인 문제를 일으킬 수 있는 신호로 해석하며(Puzakova, 2013) 서비스가 공격적이라고 인식하면 그 서비스 제공자의 도덕성을 낮게 평가할 가능성이 높다. 그러나 증진 이론가들은 사람의 특성이 변할 수 있다고 믿기 때문에 한 번의 잘못된 행동이 그 사람의 미래 행동을 예측하지 않는다고 생각한다(Puzakova, 2013). 따라서 스마트 음성 비서의 공격적 유머가 이들의 제품 태도나 도덕성 인식에 미치는 영향이 불변 이론가에 비해 덜할 것이다. 불변 이론가일수록 공격적인 유머를 더욱 부정적으로 인식하고 이를 바탕으로 서비스를 평가할 가능성이 높다. 즉, 소비자가 지닌 암묵적 이론은 스마트 음성 비서의 긍정적 및 부정적 유머가 도덕성에 미치는 영향에 차이를 만들 수 있을 것이다.

H5. 소비자의 암묵적 이론은 a) 친화적 유머, b) 자기고양적 유머, c) 자기패배적 유머, d) 공격적 유머와 지각된 도덕성 간의 관계를 조절할 것이다.

H6. 소비자의 암묵적 이론은 a) 친화적 유머, b) 자기고양적 유머, c) 자기패배적 유머, d) 공격적 유머와 소비자 태도 간의 관계를 조절할 것이다.

III. 연구 설계 및 방법

스마트 음성 비서가 사용하는 다양한 유머 스타일이 소비자의 지각된 도덕성과 태도에 미치는 영향을 살펴보기 위한 실험 자극물은 스마트 음성 비서의 유머 스타일의 4가지 유형(친화적 유머, 자기고양적 유머, 자기패배적 유

며, 공격적 유머), 그리고 유머가 없는 대화를 포함하여 5가지 시나리오로 조작하여 제시하였다.

대화 내용은 Chan et al.(2018)의 연구를 바탕으로 날씨, 자신이 잘하거나 못하는 일 등 여러 내용으로 제작하였으며, 실험참가자들은 이러한 내용을 포함한 일련의 대화 중 한 가지 시나리오만을 무작위로 제시받고 글로 작성된 대화 내용을 파악하도록 요청받았다. 다음으로 참가자들은 유머 스타일에 관한 정의 및 평가 기준(오연경, 황혜신, 2018)을 숙지한 이후, 어떤 스타일의 대화인지에 대해 평가하였다.

친화적 유머는 타인을 칭찬하고 좋은 사회적 관계를 촉진하기 위해 사용되는 유머, 자기고양적 유머는 자기 자신을 칭찬하고 스트레스나 역경 속에서도 긍정적인 마음가짐을 위해 사용되는 유머, 자기패배적 유머는 자신의 단점을 웃음 소재로 삼아 사회적 관계를 촉진하기 위해 사용되는 유머, 공격적 유머는 타인을 깎아내리고 조롱하며 배려가 없는 유머, 유머가 없는 경우란 위의 모든 유형에 속하지 않는 경우라고 제시하였다.

각 유머 스타일에 대한 평가는 앞서 언급한 조작 검사를 바탕으로 Martin et al.(2003)의 선행연구에서 정의한 5개 항목으로 7점 리커트 척도(1=전혀 그렇지 않다, 7=매우 그

렇다)를 통해 측정하였다. 지각된 도덕성은 선행연구에 따라 4개 항목으로 구성하였고(Banks, 2019), 암묵적 이론은 8개 항목으로 구성하였다(Levy, 1998). 소비자 태도는 3개 항목으로 구성하였다(Aaker & Keller, 1990). 구체적인 측정 항목은 <표 1>에 제시하였다.

자료수집은 중국 조사기업 Questionnaire Star의 유료 서비스로 진행했다. 배포된 설문지 중 불성실하거나 미완성한 것을 제외한 총 326부를 분석에 사용했다.

응답자들의 성별 분포는 남자가 167명(51.23%)으로, 친화적 유머 조건이 87건(26.68%), 자기고양적 유머 조건이 75건(23.01%), 자기패배적 유머 조건이 81건(24.85%), 공격적 유머 조건이 83건(25.46%)이었다.

스마트 음성 비서의 사용 빈도 분포는 주 3회 이상 138명(42.33%), 주 12회 121명(37.12%), 월 12회 52명(15.95%), 월 1회 이하가 15명(4.6%)으로 나타났다.

사용한 경험이 있는 스마트 음성 비서별 분포는 애플 Siri가 166명(50.92%), 샤오미 사랑이 177명(54.29%), 화웨이 샤오이 151명(46.32%), 삼성 Bixby가 132명(40.49%), 샤오도가 175명(53.68%), 티몰 요정이 149명(45.71%)으로 확인되었다.

<표 1> 확인적 요인분석 결과

요인 및 측정 문항	요인부하량(t-value)	AVE	CR
지각된 도덕성 (Cronbach's Alpha=.892)		.678	.894
이 스마트 음성 비서는 무엇이 옳고 그른 것에 대한 감각을 가지고 있다.	.793		
이 스마트 음성 비서는 도덕적 규칙에 따라 행동한다.	.881(17.603)		
이 스마트 음성 비서는 행동이 도덕적인지 아닌지 생각할 수 있다.	.810(15.866)		
이 스마트 음성 비서는 선과 악에 대해 이성적일 수 있다.	.808(15.694)		
소비자 태도 (Cronbach's Alpha=.894)		.741	.896
나는 이 스마트 음성 비서가 마음에 든다.	.840		
나는 스마트 음성 비서에 호감이 간다.	.901(19.927)		
나는 이 스마트 음성 비서가 좋다.	.840(17.886)		
암묵적 이론 (Cronbach's Alpha=.932)		.633	.932
특정한 유형의 사람이라는 것은 그 사람의 근본적인 것으로, 이는 잘 바뀌지 않는다.(R)	.811		
사람들은 다르게 행동할 수 있지만, 자신의 중요한 부분은 진정으로 변화될 수 없다.(R)	.759(15.420)		
모든 사람은 특정한 유형에 속하며, 이를 진정으로 변화시킬 수 없다.(R)	.783(16.129)		
사람들은 자신의 중요한 특성을 진정으로 변화시킬 수 없다.(R)	.784(16.146)		
사람들은 누구나 자신의 근본적인 성격을 변화시킬 수 있다.	.788(16.267)		
사람들은 자신이 어떤 유형의 사람인지를 충분히 변화시킬 수 있다.	.858(18.166)		
어떤 유형의 사람이더라도, 자신을 항상 크게 변화시킬 수 있다.	.732(14.686)		
모든 사람은 자신의 가장 근본적인 특성까지도 변화시킬 수 있다.	.842(17.636)		

IV. 실증분석

1. 조작 점검

실험 자극물 유형 간 유머 스타일의 지각 정도 차이를 점검하기 위해 일원배치 분산분석(One-way ANOVA)을 실시한 결과는 다음과 같다.

친화적 유머 대화 조건(M친화적 유머=6.00, M자기고양적 유머=3.69, M자기패배적 유머=3.65, M공격적 유머=3.89, M유머 없음=3.71; $F=28.404$, $p<0.001$), 자기고양적 유머 대화의 경우(M 친화적 유머=2.93, M 자기고양적 유머=6.08, M 자기패배적 유머=2.96, M 공격적 유머=2.94, M 유머 없음=2.76; $F=97.601$, $p<0.001$), 자기패배적 유머 대화 조건(M 친화적 유머=2.90, M 자기고양적 유머=3.28, M 자기패배적 유머=6.06, M 공격적 유머=2.81, M 유머 없음=3.10; $F=89.32$, $p<0.001$), 공격적 유머 대화의 경우(M 친화적 유머=3.10, M 자기고양적 유머=2.95, M 자기패배적 유머=3.26, M 공격적 유머=5.92, M 유머 없음=3.09; $F=78.992$, $p<0.001$), 유머가 없는 대화 조건(M 친화적 유머=3.02, M 자기고양적 유머=2.93, M 자기패배적 유머=2.93, M 공격적 유머=2.84, M 유머 없음=5.99; $F=96.342$, $p<0.001$) 등 각 조건에 맞는 유머 스타일에 대한 평가가 다른 스타일에 비해 유의하게 높게 나타나, 조작이 성공적으로 이루어졌음을 확인하였다.

2. 신뢰성 및 타당성 검증

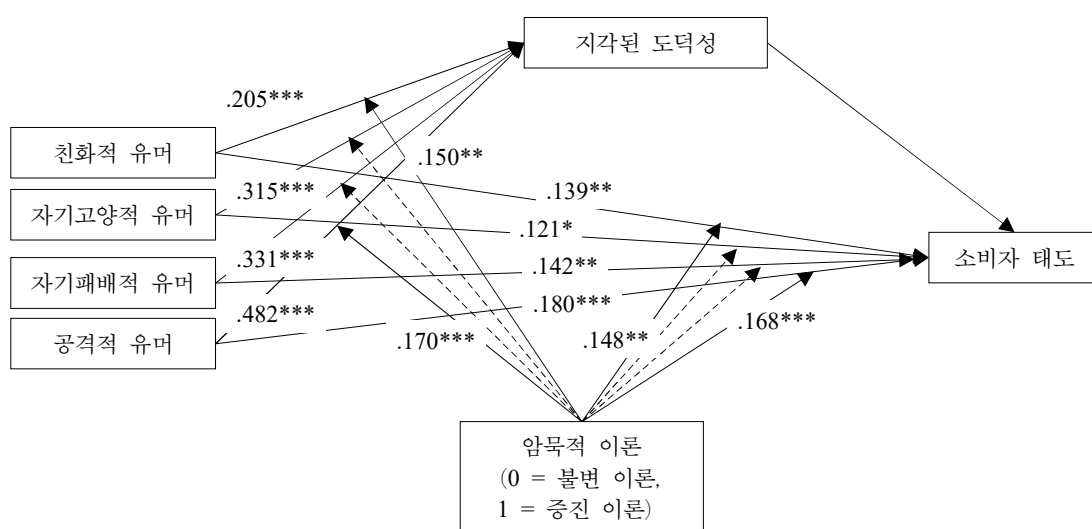
측정 항목의 타당성 및 신뢰성 검증을 위해 요인분석을 실시한 결과, 요인 적재치가 0.7 이상이었고 전체 분산의 70% 정도가 설명되어 개념 타당성을 확인하였다. Cronbach's Alpha 값이 0.8 이상으로 나타나고 신뢰성도 확보되었다. 다음으로 개별 잠재적 요인에 대한 집중타당성을 평가하기 위해 표준화 요인 부하량, 개념 신뢰도와 평균분산추출값(AVE)을 산출하였다.

측정모형의 적합도가 $\chi^2=163.313$, $\chi^2/df=1.877$, CFI=.977, TLI=.972, GFI=.940, AGFI=.917, RMSEA=.052 등 양호한 수준으로 나타났다. 각 측정변수의 표준화 요인 부하량은 0.7 이상이고 AVE의 크기는 .633~.741이며 개념신뢰도 값의 크기는 .894~.932로 나타나 구성개념의 적합성과 타당성이 입증되었다.

본 모형에서 지각된 도덕성과 암묵적 이론의 상관계수는 0.249($p<.01$)로 나타났고, 지각된 도덕성과 태도의 상관계수는 0.646($p<.01$), 태도와 암묵적 이론의 상관계수는 0.213($p<.01$)으로 나타났다. AVE 제공근의 수치는 지각된 도덕성이 0.823이고, 태도가 0.861, 암묵적 이론이 0.796으로 나타났으므로, 상관계수의 최댓값인 0.646이 AVE 제공근의 최저치인 0.796보다 낮아서 판별타당성도 확보되었다.

3. 가설검증 및 결과논의

측정 모형의 적합성 분석을 실시한 결과, 친화적 유머 모형 적합도 지수가 $\chi^2=21.571$, $\chi^2/df=1.198$, CFI=.998,



(* $p<.05$ ** $p<.01$ *** $p<.001$)

[그림 1] 연구모형과 경로분석 결과

TLI=.996, GFI=.984, AGFI=.968, RMSEA=.025 등 양호한 수준으로 나타났다. 자기고양적 유머 모형의 적합도 지수도 $\chi^2=23.435$, $\chi^2/df=1.302$, CFI=.996, TLI=.995, GFI=.982, AGFI=.964, RMSEA=0.030 등 양호한 수준으로 나타났다. 자기패배적 유머 모형 적합도 지수 또한 $\chi^2=24.052$, $\chi^2/df=1.336$, CFI=.996, TLI=.994, GFI=.982, AGFI=.964, RMSEA=.032 등 양호한 수준이었고 공격적 유머 모형 적합도 지수가 $\chi^2=21.891$, $\chi^2/df=1.216$, CFI=.998, TLI=.996, GFI=.984, AGFI=.967, RMSEA=.026 등 양호한 수준이었다.

가설에 대하여 실증분석한 검증 결과를 도식화하여 [그림 1]로 나타냈으며, <표 2>에는 세부사항을 제시하였다.

AMOS 프로그램을 활용한 경로분석 결과에 따르면, 스마트 음성 비서의 친화적 유머, 자기고양적 유머, 자기패배적 유머는 지각된 도덕성에 정(+)의 영향을 미치는 것으로 나타났다(β 친화적=.205, β 자기고양적=.315, β 자기패배적=.331, $p<.05$) 공격적 유머는 지각된 도덕성에 부(-)의 영향을 미치는 것으로 나타났다(β =-.482, $p<.05$). 그리

고 친화적 유머, 자기고양적 유머, 자기패배적 유머는 소비자 태도에 정(+)의 영향을 미치며(β 친화적=.139, β 자기고양적=.121, β 자기패배적=.142, $p<.05$) 공격적 유머는 소비자 태도에 부(-)의 영향을 미치는 것으로 나타났다(β =-.180, $p<.05$). 또한 지각된 도덕성은 소비자 태도에 정(+)의 영향을 미치는 것으로 확인되었다(β 친화적=.681, β 자기고양적=.672, β 자기패배적=.664, β 공격적=.623, $p<.05$).

다음으로, SPSS PROCESS를 활용하여 매개분석을 실시하였다. 그 결과를 살펴보면, 유머→지각된 도덕성→태도의 경로에서 모델의 간접효과 값은 친화적 유머의 경우 .140, 자기고양적 유머의 경우 .212, 자기패배적 유머의 경우 .220, 공격적 유머의 경우 -.300이고, 신뢰구간에서 0을 포함하지 않아서 각 유머와 소비자 태도 간 관계를 지각된 도덕성이 매개한다는 것을 알 수 있었다.

조절분석 결과를 살펴보면, 친화적 유머와 암묵적 이론의 상호작용항은 지각된 도덕성, 태도에 유의한 정(+)의 영향을 미치는 것으로 나타났다(β 도덕성=.150, β 태도

<표 2> 경로분석 결과

경로	표준화된 경로계수	t	p	채택여부
친화적 유머→지각된 도덕성	.205	3.541	.000	채택
친화적 유머→태도	.139	3.048	.002	채택
지각된 도덕성→태도	.681	11.295	.000	채택
친화적 유머→지각된 도덕성→태도	.140		.001	채택
친화적 유머*암묵적 이론→지각된 도덕성	.150	3.038	.004	채택
친화적 유머*암묵적 이론→태도	.148	2.904	.004	채택
자기고양적 유머→지각된 도덕성	.315	5.506	.000	채택
자기고양적 유머→태도	.121	2.546	.011	채택
지각된 도덕성→태도	.672	10.908	.000	채택
자기고양적 유머→지각된 도덕성→태도	.212		.000	채택
자기고양적 유머*암묵적 이론→지각된 도덕성	-.087	-1.674	.095	기각
자기고양적 유머*암묵적 이론→태도	-.067	-1.287	.199	기각
자기패배적 유머→지각된 도덕성	.331	5.767	.000	채택
자기패배적 유머→태도	.142	2.976	.003	채택
지각된 도덕성→태도	.664	10.914	.000	채택
자기패배적 유머→지각된 도덕성→태도	.220		.000	채택
자기패배적 유머*암묵적 이론→지각된 도덕성	-.081	-1.585	.114	기각
자기패배적 유머*암묵적 이론→태도	-.071	-1.403	.161	기각
공격적 유머→지각된 도덕성	-.482	-8.637	.000	채택
공격적 유머→태도	-.180	-3.483	.000	채택
지각된 도덕성→태도	.623	9.911	.000	채택
공격적 유머→지각된 도덕성→태도	-.300		.000	채택
공격적 유머*암묵적 이론→지각된 도덕성	.170	6.458	.000	채택
공격적 유머*암묵적 이론→태도	.168	6.937	.000	채택

$=.148, p<.05$). 따라서 암묵적 이론이 친화적 유머와 지각된 도덕성 또는 태도 간의 관계를 긍정적으로 조절하는 효과가 있음을 알 수 있다[그림 2].

공격적 유머와 암묵적 이론의 상호작용항은 지각된 도덕성, 태도에 유의한 정(+)의 영향을 미치는 것으로 나타났다(β 도덕성 $=.170$, β 태도 $=.168$, $p<.05$). 따라서 암묵적 이론이 공격적 유머와 지각된 도덕성 또는 태도 간의 관계를 긍정적으로 조절하는 효과가 나타났다[그림 3]. 자기고양적 유머나 자기패배적 유머와 암묵적 이론의 상호작용항은 유의하지 않은 것으로 나타났다.

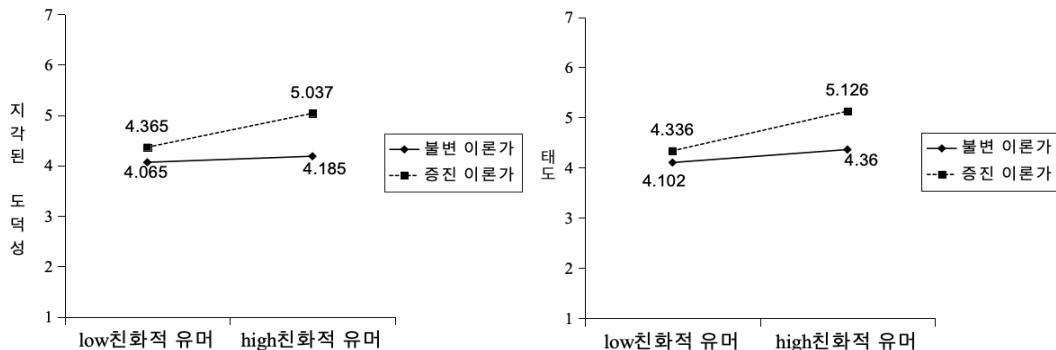
V. 결론

스마트 음성 비서의 언어 스타일은 사용자 경험에 큰 영향을 미친다. 개발자들은 사용자와의 소통을 더욱 즐겁게 만들기 위해 유머러스한 요소를 포함시키기도 한다. 본 연구에서는 스마트 음성 비서가 사용하는 네 가지 유머 스타일(친화적 유머, 자기고양적 유머, 자기패배적 유머, 공격적 유머)이 스마트 음성 비서에 대해 소비자들이 지각하는

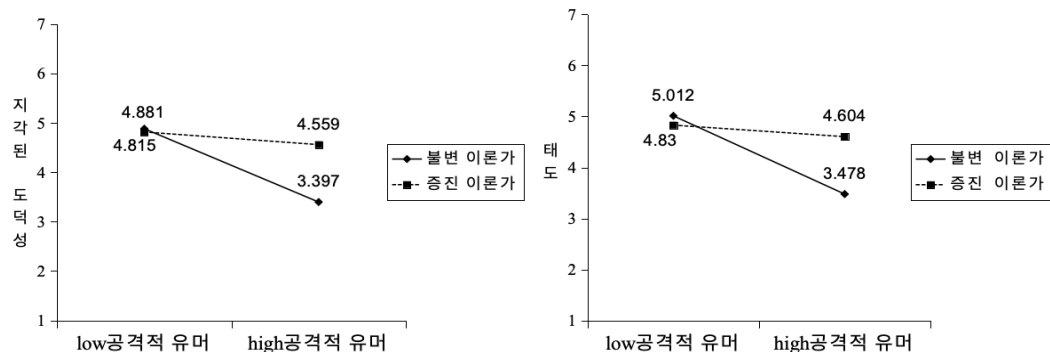
도덕성과 태도에 미치는 영향을 확인하고자 하였다. 또한, 지각된 도덕성이 스마트 음성 비서의 유머 스타일과 소비자 태도 사이에서 매개 역할을 하는지 조사하였다. 마지막으로, 소비자들이 가진 암묵적 이론이 이 과정을 조절하는 효과를 나타내는지를 탐구했다. 연구 결과는 다음과 같다.

첫째, 스마트 음성 비서는 친화적 유머, 자기고양적 유머, 자기패배적 유머를 사용할 때 지각된 도덕성과 소비자 태도에 긍정적인 영향을 미치지만, 공격적 유머의 사용은 지각된 도덕성과 소비자 태도에 부정적인 영향을 미치는 것으로 나타났다. 공격적 유머는 일부 소비자에게 흥미롭게 여겨질 수도 있겠지만 스마트 음성 비서에 대한 지각된 도덕성을 높이는 데 있어 다른 유머 스타일만큼 효과적이지 않다는 것을 보여준다. 이는 고객 불만 해결의 맥락에서 공격적 유머보다 친화적 유머를 사용할 것을 주장한 Shin과 Larson(2020)의 연구와도 일치하며, 이러한 발견은 소비자 행동에 대한 유머 효과의 미묘함과 복잡성을 이해하는 것이 중요함을 강조한다(Mathies et al., 2016; Schanke et al., 2021).

둘째, 소비자 지각된 도덕성이 소비자 태도에 긍정적인 영향을 미쳤다. 또 스마트 음성 비서가 다양한 유머 스타일



[그림 2] 친화적 유머와 암묵적 이론의 상호작용 효과



[그림 3] 공격적 유머와 암묵적 이론의 상호작용 효과

을 사용하여 태도에 미치는 영향은 지각된 도덕성을 통해 매개된다는 것을 확인하였다. 따라서 스마트 음성 비서가 친화적 유머, 자기고양적 유머, 자기패배적 유머 스타일을 많이 사용하면 소비자들은 이 비서가 더 도덕적이라고 느낄 수 있고, 스마트 음성 비서에 대한 태도도 긍정적으로 변화할 가능성이 높다. 반면, 공격적 유머를 사용할 때는 도덕성 인식 수준이 낮아질 수 있으므로 신중한 고려가 필요함을 시사한다.

로봇의 행동은 인간의 행동과 유사한 지각 효과를 유발하는데, 이는 인간의 뇌가 로봇이 과제를 수행하는 방식을 인간이 동일한 과제를 수행하는 것과 유사하게 이해할 가능성이 높다는 것을 의미한다(Wykowska et al., 2014). AI가 행동을 할 때 고의적이라면 사람들은 그 행동의 영향이나 더 광범위한 결과를 인식할 능력이 없을 수 있다. 인식하고 의도하는 정신력이 부족하다는 것은 AI가 비슷한 상황에서 인간보다 더 무책임하다는 것을 의미한다(Wilson et al., 2022). 따라서 로봇이나 인공지능 기술의 역할이 커질수록 AI 시스템을 어떻게 설계하고 사용할 것인지에 대한 중요한 윤리적 문제가 제기된다. 인공지능은 의도치 않게 인간 사회에 해를 끼칠 수 있으며, 예를 들어 특정 알고리즘을 통한 편견이 그 예이다(Owe & Baum, 2021). 그러므로 스마트 음성 비서와 같은 인공지능 기술은 사람들에게 광범위한 도덕적 관심과 인식을 요구하게 된다.

셋째, 암묵적 이론은 친화적 유머와 공격적 유머가 지각된 도덕성 또는 태도 간의 관계를 긍정적으로 조절하는 효과가 있다. 발전 가능성이 있다고 믿는 증진 이론가일수록 비서가 사용하는 친화적 유머의 수준이 지각된 도덕성과 태도에 미치는 긍정적인 영향력이 증가하는 경향이 있다. 반면, 사람은 잘 변하지 않는다고 생각하는 불변 이론가들은 비서가 사용하는 공격적 유머의 수준이 지각된 도덕성과 태도에 미치는 부정적인 영향력이 증가하는 결과가 나타났다. 이러한 결론은 불변 이론가의 경우 부정적 행동의 책임을 행위자의 본래 특성으로 돌리는 반면, 증진 이론가는 다른 상황에서 행동이 다르게 나타날 수 있음을 허용한다(Puzakova, 2013)는 점과도 맥락을 같이 한다.

이상의 연구 결과를 통해 다음과 같은 시사점을 제안할 수 있다.

시스템 개발자들은 스마트 음성 비서에 유머 기능을 추가할 때 다양한 유머 스타일을 신중하게 고려해야 한다. 또한 소비자의 암묵적 이론 차이를 반영하여 유머 스타일을 적절히 조절할 필요가 있다. 친화적 유머를 사용하는 스마트 음성 비서는 소비자에게 쉽고 즐거운 대화 분위기를 제

공하며, 대화의 긴장감을 줄여준다. 자기고양적 유머를 사용하는 스마트 음성 비서는 긍정적이고 활기찬 분위기를 조성하여 긴장, 우울, 불안감을 완화시킨다. 자기패배적 유머를 사용하는 스마트 음성 비서는 겸손한 태도를 보여주어 자연스러운 대화를 이끌어낸다. 반면, 공격적 유머는 지나치게 공격적일 경우 스마트 음성 비서를 무례하고 비도덕적으로 보이게 하여 결국 소비자의 태도를 악화시킬 수도 있다.

개인차의 암묵적 이론이 조절 효과에 미치는 영향에도 주목해야 한다. 증진 이론가는 친화적 유머를 긍정적으로 받아들이고, 스마트 음성 비서가 더 나은 방향으로 발전할 것이라는 기대를 할 수 있다. 반면, 불변 이론가는 공격적 유머 수준이 높아질수록 부정적인 반응을 보일 가능성이 크며, 공격적 유머가 스마트 음성 비서 자체의 단점을 대표한다고 인식하거나 비판적인 태도를 형성할 가능성이 높다. 이에 여러 소비자 집단이 보이는 차별적 반응을 살필 필요가 있을 것이다.

본 연구는 몇 가지 한계점을 제시하며 미래 연구 방향을 다음과 같이 제안한다. 첫째, 스마트 음성 비서는 현재 마케팅에서 대표적인 제품 중 하나이지만, 관광, 호텔, 금융 등 다양한 분야에서 많은 관심을 가진다. 향후 연구는 산업에 적용한 결과를 분석하여 연구의 적용 범위를 넓힐 수 있을 것이다.

둘째, 향후 연구에서는 젊은 층과 노인층을 대상으로 더 구체적인 연구를 진행함으로써 각 그룹에 더 적합한 스마트 음성 비서를 개발할 수 있을 것이다.

셋째, 스마트 음성 비서의 공격적인 유머를 수용하거나 심지어 좋아하는 소비자도 있다. 이러한 다양한 상황을 반영하고 소비자의 반응을 탐색할 필요가 있다.

주제어: 소비자 태도, 스마트 음성 비서, 유머 스타일, 암묵적 이론, 지각된 도덕성

REFERENCES

- 모은가, 류미현(2023). 중국 온라인 쇼핑몰의 챗봇 서비스에 대한 혜택, 희생, 사회적 실재감이 지각된 가치 및 지속이용의도에 미치는 영향 연구. *한국생활과학회지*, 32(4), 409-423.
- 오연정, 황혜신(2018). 부모의 유머스타일이 유아의 정서

- 표현성에 미치는 영향. *한국보육지원학회지*, 14(6), 107-122.
- Aaker, D. A., & Keller, K. L. (1990). Consumer evaluations of brand extensions. *Journal of marketing*, 54(1), 27-41.
- Banks, J. (2019). A perceived moral agency scale: Development and validation of a metric for humans and social machines. *Computers in Human Behavior*, 90, 363-371.
- Batool, S., Panezai, B. A., Baloch, G. J., & Sohail, S. (2023). Influence of customer perceived ethicality on customer loyalty in the banking industry of Pakistan. *International Journal of Ethics and Systems*, 39(1), 54-80.
- Beard, F. K. (2008). Advertising and audience offense: The role of intentional humor. *Journal of Marketing Communications*, 14(1), 1-17.
- Behera, R. K., Bala, P. K., & Ray, A. (2024). Cognitive chatbot for personalised contextual customer service: Behind the scene and beyond the hype. *Information Systems Frontiers*, 26(3), 899-919.
- Braslavski, P., Blinov, V., Bolotova, V., & Pertsova, K. (2018). How to evaluate humorous response generation, seriously? In *Proceedings of the 2018 Conference on Human Information Interaction & Retrieval*, 225-228.
- Brunk, K. H. (2012). Unethical company and brand perceptions: Conceptualising and operationalising consumer meanings. *Journal of Business Ethics*, 111(4), 551 - 565.
- Campbell, L., Martin, R. A., & Ward, J. R. (2008). An observational study of humor use while resolving conflict in dating couples. *Personal Relationships*, 15(1), 41-55.
- Chan, Y. C., Hsu, W. C., Liao, Y. J., Chen, H. C., Tu, C. H., & Wu, C. L. (2018). Appreciation of different styles of humor: An fMRI study. *Scientific Reports*, 8(1), 15649.
- Chaves, A. P., & Gerosa, M. A. (2021). How should my chatbot interact? A survey on social characteristics in human-chatbot interaction design. *International Journal of Human-Computer Interaction*, 37(8), 729-758.
- Chiu, C. Y., Hong, Y. Y., & Dweck, C. S. (1997). Lay dispositionism and implicit theories of personality. *Journal of Personality and Social Psychology*, 73(1), 19.
- Chung, H., & Zhao, X. (2003). Humour effect on memory and attitude: moderating role of product involvement. *International Journal of Advertising*, 22(1), 117-144.
- Chung, M., Ko, E., Joung, H., Kim, S. J. (2020). Chatbot e-service and customer satisfaction regarding luxury brands. *Journal of Business Research*, 117, 587-595.
- Croes, E. A., & Antheunis, M. L. (2021). Can we be friends with Mitsuku? A longitudinal study on the process of relationship formation between humans and a social chatbot. *Journal of Social and Personal Relationships*, 38(1), 279-300.
- Dang, J., & Liu, L. (2022). A growth mindset about human minds promotes positive responses to intelligent technology. *Cognition*, 220, 104985.
- Duncan, W. J., Feisal, J. P. (1989). No laughing matter: Patterns of humor in the workplace. *Organizational Dynamics*, 17(4), 18 - 31.
- Eisend, M. (2009). A meta-analysis of humor in advertising. *Journal of the Academy of Marketing Science*, 37(2), 191 - 203.
- Ford, T. E., Wentzel, E. R., & Lorion, J. (2001). Effects of exposure to sexist humor on perceptions of normative tolerance of sexism. *European Journal of Social Psychology*, 31(6), 677-91.
- Förster, K., & Brantner, C. (2016). Masking the offense? An ethical view on humor in advertising. *Journal of Media Ethics*, 31(3), 146-161.
- Følstad, A., Nordheim, C. B., & rkli, C. A. (2018). What makes users trust a chatbot for customer service? An exploratory interview study. In *Internet Science: 5th International Conference, INSCI 2018, St. Petersburg, Russia, October 24-26, 2018, Proceedings* 5, 194-208. Springer International Publishing.
- Grudin, J., & Jacques, R. (2019). Chatbots, humbots, and the quest for artificial general intelligence. In

- Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, 1-11.
- Hong, Y. Y., Chiu, C. Y., Dweck, C. S., Lin, D. M. S., & Wan, W. (1999). Implicit theories, attributions, and coping: a meaning system approach. *Journal of Personality and Social Psychology*, 77(3), 588.
- Jiang, T., Li, H., & Hou, Y. (2019). Cultural differences in humor perception, usage, and implications. *Frontiers in Psychology*, 10, 123.
- Kellaris, J. J., Cline, T. W. (2007). Humor and ad memorability: On the contributions of humor expectancy, relevancy, and need for humor. *Psychology & Marketing*, 24(6), 497 - 509.
- Kwon, J., & Nayakankuppam, D. (2015). Strength without elaboration: The role of implicit self-theories in forming and accessing attitudes. *Journal of Consumer Research*, 42(2), 316-339.
- Levy, S. R., Stroessner, S. J., & Dweck, C. S. (1998). Stereotype formation and endorsement: The role of implicit theories. *Journal of Personality and Social Psychology*, 74(6), 1421.
- Liu, D., Lv, Y., & Huang, W. (2023). How do consumers react to chatbots' humorous emojis in service failures. *Technology in Society*, 73, 102244.
- Lopatovska, I. (2020). Classification of humorous interactions with intelligent personal assistants. *Journal of Librarianship and Information Science*, 52(3), 931-942.
- Markovic, S., Iglesias, O., Singh, J. J., & Sierra, V. (2018). How does the perceived ethicality of corporate services brands influence loyalty and positive word-of-mouth? Analyzing the roles of empathy, affective commitment, and perceived quality. *Journal of Business Ethics*, 148, 721-740.
- Martin, R. A., Puhlik-Doris, P., Larsen, G., Gray, J., Weir, K. (2003). Individual differences in uses of humor and their relation to psychological well-being: Development of the humor styles questionnaire. *Journal of Research in Personality*, 37(1), 48-75.
- Mathies, C., Chiew, T. M., & Kleinaltenkamp, M. (2016). The antecedents and consequences of humour for service: A review and directions for research. *Journal of Service Theory and Practice*, 26(2), 137-162.
- McColl-Kennedy, J. R., & Sparks, B. A. (2003). Application of fairness theory to service failures and service recovery. *Journal of Service Research*, 5(3), 251-266.
- Niculescu, A., Van Dijk, B., Nijholt, A., Li, H., & See, S. L. (2013). Making social robots more attractive: the effects of voice pitch, humor and empathy. *International Journal of Social Robotics*, 5, 171-191.
- Norrick, N. R., Spitz, A. (2008). Humor as a resource for mitigating conflict in interaction. *Journal of Pragmatics*, 40(10), 1661-1686.
- Owe, A., & Baum, S. D. (2021). Moral consideration of nonhumans in the ethics of artificial intelligence. *AI and Ethics*, 1(4), 517-528.
- Puzakova, M., Kwak, H., & Rocereto, J. F. (2013). When humanizing brands goes wrong: The detrimental effect of brand anthropomorphization amid product wrongdoings. *Journal of Marketing*, 77(3), 81-100.
- Schanke, S., Burtch, G., Ray, G. (2021). Estimating the impact of "humanizing" customer service chatbots. *Information Systems Research*, 32(3), 736 - 751.
- Scott, T. (2007). Expression of humour by emergency personnel involved in sudden deathwork. *Mortality*, 12(4), 350-364.
- Sheehan, B., Jin, H. S., Gottlieb, U. (2020). Customer service chatbots: Anthropomorphism and adoption. *Journal of Business Research*, 115, 14-24.
- Shin, H., & Larson, L. R. (2020). The bright and dark sides of humorous response to online customer complaint. *European Journal of Marketing*, 54(8), 2013-2047.
- Shin, H., Bunosso, I., & Levine, L. R. (2023). The influence of chatbot humour on consumer evaluations of services. *International Journal of Consumer Studies*, 47(2), 545-562.
- Sierra, V., Iglesias, O., Markovic, S., & Singh, J. J. (2017). Does ethical image build equity in corporate services brands? The influence of

- customer perceived ethicality on affect, perceived quality, and equity. *Journal of Business Ethics*, 144, 661-676.
- Singh, J. J., Iglesias, O., & Batista-Foguet, J. M. (2012). Does having an ethical brand matter? The influence of consumer perceived ethicality on trust, affect and loyalty. *Journal of Business Ethics*, 111, 541-549.
- Song, D., & Rim, H. (2024). Understanding the influence of implicit theories on donor's preference for time or money. *Journal of Consumer Behaviour*.
- Thomaz, F., Salge, C., Karahanna, E., Hulland, J. (2020). Learning from the dark web: Leveraging conversational agents in the era of hyper-privacy to enhance marketing. *Journal of the Academy of Marketing Science*, 48(1), 43-63.
- Thorson, J. A., Powell, F. C., Sarmany-Schuller, I., & Hampes, W. P. (1997). Psychological health and sense of humor. *Journal of Clinical Psychology*, 53(6), 605-619.
- Toader, D. C., Boca, G., Toader, R., Măcelaru, M., Toader, C., Ighian, D., Rădulescu, A. T. (2020). The effect of social presence and chatbot errors on trust. *Sustainability*, 12(1), 256.
- Van Kuilenburg, P., de Jong, M. D., & Van Rompay, T. J. (2011). 'That was funny, but what was the brand again?' Humorous television commercials and brand linkage. *International Journal of Advertising*, 30(5), 795-814.
- Warren, C., & McGraw, A. P. (2016). When does humorous marketing hurt brands?. *Journal of Marketing Behavior*, 2(1), 39-67.
- Warren, C., Barsky, A., McGraw, A. P. (2018). Humor, comedy, and consumer behavior. *Journal of Consumer Research*, 45(3), 529 - 552.
- Wilson, A., Stefanik, C., & Shank, D. B. (2022). How do people judge the immorality of artificial intelligence versus humans committing moral wrongs in real-world situations? *Computers in Human Behavior Reports*, 8, 100229.
- Wykowska, A., Chellali, R., Al-Amin, M. M., & Muller, H. J. (2014). Implications of robot actions for human perception. How do we represent actions of the observed robots? *International Journal of Social Robotics*, 6, 357-366.
- Yang, Z., Zhou, J., & Yang, H. (2023). The impact of AI's response method on service recovery satisfaction in the context of service failure. *Sustainability*, 15(4), 3294.
- Yue, X. D., Liu, K. W. Y., Jiang, F., & Hiranandani, N. A. (2014). Humor styles, self-esteem, and subjective happiness. *Psychological Reports*, 115(2), 517-525.

Received 06 November 2024;

Accepted 27 November 2024