

신뢰 지향 아바타 챗봇의 설계와 사용자 인식: 전공 선택과 진로 상담 사례

성 봉 주¹ · 이 규 성¹ · 김 정 환² · 박 대 근^{2*}

¹차의과학대학교 일반대학원 AI헬스케어융합학과 석사과정

²차의과학대학교 일반대학원 AI헬스케어융합학과 교수

Trust-by-Design of an Avatar Chatbot: Advising University Students on Their Majors

Bong Ju Sung¹ · Gyu Seong Lee¹ · Junghwan Kim² · Dae Keun Park^{2*}

¹M.S. Candidate, Department of AI Healthcare Convergence, CHA University, Pocheon 11160, Korea

²Professor, Department of AI Healthcare Convergence, CHA University, Pocheon 11160, Korea

[요 약]

AI 챗봇의 사용자 수용은 시스템 신뢰에 좌우되며, 검색 증강 생성(RAG)과 아바타 인터페이스가 신뢰 형성의 핵심 방법으로 연구되어 왔다. 그러나 기존 연구는 개별 신뢰 요소를 단편적으로 검증해 왔을 뿐, 다층의 신뢰 신호를 단일 시스템에 통합 설계하는 신뢰 지향(Trust-by-Design) 접근의 실증 사례는 부족하였다. 본 연구는 신뢰 이론을 종합하여 챗봇 정체성·답변 품질·대화 자연스러움·운영 정책 투명성에 걸친 신뢰 신호를 도출하고, 멀티모달 아바타 챗봇에 반영하여 대학생 전공 선택과 진로 상담에 적용하였다. 운영 정책 신호의 인식 수준이 챗봇 정체성 신호보다 높게 나타났고 답변 품질 내부에서도 정서적 측면이 인지적 측면보다 약하게 전달되었으며, 격차의 원인은 기술적 한계, 정보 깊이 부족, 분위기 불일치, 발화 주체에 대한 통합 인식 부족으로 해석된다. 본 연구는 신뢰 지향 챗봇을 구현하고 사용자 인식 차원에서 검증한 결과를 제시함으로써 후속 연구의 기반을 마련하였다.

[Abstract]

User acceptance of AI chatbots depends on user trust, and retrieval-augmented generation and avatar interfaces have been studied as core methods of trust formation. Prior studies, however, have largely examined trust elements in isolation, leaving few empirical cases of trust-by-design approaches that integrate multi-layered trust signals into a single system. This study derives trust signals spanning chatbot identity, response quality, interaction naturalness, and operational-policy transparency, and applies them to a multimodal avatar chatbot intended to advise university students on their majors. Operational-policy signals were perceived more strongly than chatbot-identity signals; affective aspects of response quality were transmitted less strongly than cognitive aspects; and the gap reflects technical limitations, insufficient information depth, persona-tone mismatch, and incomplete integration of speaking-agent cues. By implementing the trust-by-design chatbot and examining its reach to user perception, this study provides a foundation for future research.

색인어 : 신뢰 지향 설계, 멀티모달 인터페이스, 아바타 챗봇, 검색 증강 생성, 전공 선택·진로 상담

Keyword : Trust-by-Design, Multimodal Interface, Avatar Chatbot, Retrieval-Augmented Generation, Major Advising

<http://dx.doi.org/10.9728/dcs.2026.27.7.1923>



This is an Open Access article distributed under the terms of the Creative Commons Attribution Non-Commercial License (<http://creativecommons.org/licenses/by-nc/3.0/>) which permits unrestricted non-commercial use, distribution, and reproduction in any medium, provided the original work is properly cited.

Received 07 May 2026; **Revised** 08 June 2026

Accepted 06 July 2026

***Corresponding Author; Dae Keun Park**

Tel: +82-31-850-8984

E-mail: dkpark@cha.ac.kr

1. 서 론

1-1 연구배경

AI 챗봇은 학습·상담·고객 응대 등 다양한 도메인에 빠르게 도입되고 있으며, 대규모 언어 모델(Large Language Model, LLM)의 발전에 힘입어 자연스러운 대화 능력을 확보하기에 이르렀다. 그러나 챗봇의 기술적 성능 향상이 곧 사용자 수용으로 이어지지 않는다는 사실은 사용자 챗봇을 신뢰하지 않으면 응답을 받아들이지 않으며, 신뢰가 결여된 시스템은 지속 사용으로 연결되지 않는다[1],[2]. 따라서 챗봇의 사용자 수용을 결정짓는 핵심 변수는 시스템의 객관적 성능이 아니라 사용자가 그 시스템에 부여하는 신뢰이며, 신뢰 형성에 관한 이해는 챗봇 설계의 본질적 과제로 부상하고 있다[3].

신뢰 형성에 관한 최근 연구는 크게 두 흐름으로 정리된다. 첫째, 답변의 사실성을 확보하기 위한 검색 증강 생성(Retrieval-Augmented Generation, RAG)이다. RAG는 LLM의 환각(hallucination) 위험을 완화하고 도메인 지식의 정확성을 담보하는 기술적 방법으로 자리잡았으며[4]~[6], 교육·상담 도메인에서 신뢰 형성의 인지적 기반을 제공한다. 둘째, 사용자에게 사회적 실재감(social presence)을 부여하는 아바타 인터페이스이다. 시각·음성·표정이 통합된 아바타는 텍스트 단독 인터페이스에 비해 풍부한 사회적 단서를 제공함으로써 의인화(anthropomorphism)와 라포(rapport)를 매개로 신뢰 형성에 기여하는 것으로 보고되어 왔다[7]~[9].

그러나 두 방법이 개별적으로 신뢰 형성에 기여한다는 점이 곧 이를 결합한 시스템이 충분한 신뢰를 형성한다는 것을 의미하지는 않는다. 신뢰는 단일 기능으로 만들어지는 단순 변수가 아니라, 챗봇 정체성·답변 품질·대화 자연스러움·운영 정책 투명성 등 다층의 신호가 사용자에게 통합적으로 도달할 때 비로소 형성되는 복합 구성개념이다[10]~[12]. 따라서 신뢰 형성을 시스템 설계의 사후 평가 대상으로 다루는 접근에서 벗어나, 설계 단계에서 다층 신호를 의도적으로 구조화하여 내재화하는 신뢰 지향(Trust-by-Design) 접근이 요청되며, 그 설계 의도가 사용자 인식에 실제로 도달하는지를 함께 검증하는 작업이 동반되어야 한다.

본 연구는 이러한 접근의 실증 사례로 C대학 신입생을 대상으로 한 전공 선택·진로 상담 챗봇을 다룬다. C대학은 1학년 전체를 무전공으로 선발하여 2학년 진학 시 전공을 결정하도록 운영하고 있어, 신입생은 입학 직후부터 전공 정보 탐색과 진로 의사결정에 대한 높은 수요를 갖는다.

1-2 연구 문제와 연구 질문

선행 연구는 RAG의 사실성 효과[4],[6], 아바타의 사회적 실재감 효과[7],[8], 챗봇 신뢰의 다차원성[10],[13] 등 개별 요소의 신뢰 효과를 단편적으로 검증해 왔다. 그러나 이러

한 다층의 신뢰 신호를 단일 시스템에 통합 설계하고, 그 설계 의도가 실제 사용자에게 어떻게 전달되는가를 실증적으로 분석한 사례는 충분히 축적되지 않았다. 특히 한국 대학 환경에서 학생을 대상으로 한 신뢰 지향 챗봇의 설계와 평가 사례는 매우 드물다.

이에 본 연구는 두 가지 연구 질문을 설정한다. RQ1은 시스템 구현에 관한 것이며, RQ2는 그 구현이 사용자 인식에 도달하는 양상에 관한 것으로, 본 연구는 두 질문을 함께 검증한다.

- RQ1: 신뢰 이론을 통합한 신뢰 지향 멀티모달 아바타 챗봇을 설계·구현할 수 있는가?
- RQ2: 18개 신뢰 지향 설계 요소(F1~F18)가 사용자에게 어떻게 도달하는가?

RQ2는 탐색적 횡단 설문을 통한 기술 통계 분석에 한정되며, 변수 간 인과 관계 추론은 후속 연구의 과제로 남긴다. 본 연구의 사례 도메인은 대학생의 전공 선택과 진로 상담으로 설정하였다. 학과 사무실 운영 시간의 제약과 학사 정보의 분산은 학생의 정보 접근 비용을 높이는 대표적 환경이며, 신뢰 지향 챗봇의 적용 가치가 명확히 드러나는 영역이기 때문이다.

1-3 연구의 기여 및 논문 구성

본 연구는 신뢰 지향 설계를 멀티모달 아바타 챗봇에 구현하고 사용자 인식 차원에서 검증함으로써, 분석 체계·시스템 아키텍처·실증 사례의 세 측면에서 기여한다. 첫째, 신뢰 이론을 종합하여 챗봇 정체성·답변 품질·대화 자연스러움·운영 정책 투명성에 걸친 18개 신뢰 컴포넌트로 구성된 신뢰 지향 설계 체계를 제안한다. 둘째, 단일 백엔드에서 아바타 대면형(FTF), 음성 기반(STS), 텍스트 기반(TTT)의 세 모달리티를 통합 운영하는 시스템 아키텍처의 설계와 구현을 보고한다. 셋째, 국내 대학 신입생의 전공 선택·진로 상담 맥락에서 멀티모달 신뢰 지향 챗봇을 설계·실증한 사례를 보고하여, 신뢰 형성 검증을 위한 후속 연구의 분석 틀을 제공한다.

본 논문은 이론적 배경에서 출발하여 시스템 설계, 실증 결과, 논의의 순서로 신뢰 지향 설계의 구현과 사용자 인식 도달성을 차례로 검증한다. 2장은 신뢰 이론, RAG, 아바타 인터페이스에 관한 선행 연구를 정리한다. 3장은 시스템 설계와 18 컴포넌트 체계, 설문 설계를 기술한다. 4장은 C대학 신입생 50명을 대상으로 한 설문 결과를 컴포넌트별 인지율과 영역별 평균 비교로 나누어 보고한다. 5장은 결과에서 관찰된 인지율 격차의 원인을 자유응답 분석에 근거해 논의하고, 본 연구의 함의·한계·후속 연구 방향을 제시한다.

II. 선행 연구

2-1 AI 챗봇 신뢰 이론

신뢰는 상대방이 자신에게 이로운 방식으로 행동할 것이라는 기대를 기반으로 취약성을 수용하는 심리적 상태로 정의된다[12]. AI 챗봇 맥락에서 신뢰는 사용자가 시스템의 역량과 의도를 긍정적으로 평가하고 그 출력에 의존하려는 의향으로 개념화되며, 단일 차원이 아닌 인지적·정서적·행동적 차원이 결합된 복합 구성개념으로 이해된다.

Mayer 외[12]가 제안한 ABI 모형은 신뢰의 구성 요소를 능력(Ability), 선의(Benevolence), 진실성(Integrity)의 세 차원으로 정식화하였다. AI 챗봇에 적용하면 능력은 정확하고 맥락에 적합한 정보를 제공하는 시스템 성능, 선의는 사용자 이익 중심의 응답 설계, 진실성은 일관된 발화 원칙과 불확실성의 명시적 고지로 각각 대응된다. Hoff와 Bashir[14]는 자동화 시스템에 대한 신뢰를 기질적·상황적·학습된 신뢰의 3계층으로 분류하였고, McKnight 외[15]는 초기 신뢰(initial trust) 형성의 과정을 정식화하여 사전 경험이 부재한 신규 사용자의 신뢰 형성을 설명하였다.

국내 실증 연구는 챗봇 신뢰의 다차원적 형성을 일관되게 보고해 왔다. 정훈실과 김영인[10]은 패션 챗봇의 상호작용·정보·시스템·디자인 품질이 챗봇 신뢰를 통해 브랜드 신뢰로 전이되는 경로를 보고하였으며, 이유리 외[13]는 정보시스템 성공모형의 관점에서 정보·시스템·서비스 품질이 신뢰를 거쳐 사용의도에 작용함을 검증하였다. 정훈실과 김영인[10]은 패션 챗봇의...전이되는 경로를 보고하였으며, 이유리 외[13]는 정보시스템성공모형의 관점에서 정보·시스템·서비스 품질이 신뢰를 거쳐 사용의도에 작용함을 검증하였다. 임 외[11]는 대학생의 생성형 AI 학습 활용 맥락에서 챗봇 서비스 품질·신뢰·만족·지속사용의도의 관계를 실증하여, 본 연구의 신입생 대상 챗봇 평가와 직접 맞는 선행 근거를 제공한다. 김민성과 구철모[16]는 ChatGPT 사용 맥락에서 신뢰가 인지적 신뢰와 정서적 신뢰의 두 차원으로 분화되어 행동의도에 차별적으로 작용함을 확인하였다. 이러한 분화는 본 연구의 신뢰 지향 설계가 사실 정확성(인지적)과 인터페이스의 따뜻함·자연스러움(정서적)을 모두 포괄해야 함을 시사한다.

미디어 이론 측면에서 Nass 외[17]가 제안한 CASA(Computers Are Social Actors) 패러다임은 사람들이 컴퓨터를 사회적 행위자로 지각함을 입증하였고, Lombard와 Xu[18]는 이를 MASA(Media Are Social Actors) 패러다임으로 확장하였다. 박선영과 이상원[19]은 자기 노출과 의인화 전략이 친밀감을 매개로 신뢰에 영향을 미친다는 것을 실험적으로 보였다. Lee와 See[20]는 자동화 신뢰가 적정 수준(calibrated trust)으로 조정되어야 함을 강조하였으며, Jacovi 외[21]는 AI 시스템에 대한 신뢰의 형식적 조건과 목표를 정식화하였다.

2-2 검색 증강 생성과 환각 완화

LLM은 광범위한 언어 이해·생성 능력을 지니지만, 학습 데이터에 포함되지 않은 도메인 특화 정보나 최신 정보를 요구하는 질의에 대응할 때 사실 근거 없는 응답, 즉 환각이 빈번히 발생한다[22]. 학과 커리큘럼·진로 정보·자격증 등 정확도 요구가 높고 갱신 주기가 짧은 상담 도메인에서는 환각이 시스템 신뢰성을 심각하게 위협한다. RAG[23]는 LLM의 파라미터 외부에 구축된 지식베이스에서 관련 문서를 검색하고 이를 프롬프트에 결합한 뒤 응답을 생성하는 구조로, 모델이 실제 지식베이스 내 문서에 근거하여 응답하도록 유도함으로써 환각을 억제한다.

교육 분야에서 RAG의 응용은 빠르게 확산되고 있다. Swacha와 Gracel[4]은 교육용 RAG 챗봇 응용 사례를 종합하였고, Li 외[5]는 청크 분할 전략과 검색 임계값 설정이 응답 품질에 결정적 영향을 미침을 보고하였다. 환각 완화 방법 측면에서 Zhang과 Zhang[24]은 RAG 기반 모델의 환각 완화 방법론을 체계적으로 검토하였다. RAG 시스템의 검색 단계에서는 의미 검색이 핵심 역할을 담당하며, 다언어·다기능·다단위 임베딩을 지원하는 모델[25] 등의 고성능 모델이 활용된다. 한편 Liao와 Sundar[26]는 신뢰 형성을 의사소통 관점에서 재구성하여 AI 시스템이 사용자에게 보내는 신호의 명시성과 일관성이 책임 있는 신뢰의 핵심 조건임을 제시하였다.

2-3 아바타 인터페이스와 사회적 실재감

대화형 에이전트는 인간의 대화적 행동을 모사하는 가상 캐릭터 또는 아바타로, 텍스트 또는 음성만으로 구성된 인터페이스에 비해 사용자와의 상호작용에 사회적 맥락을 부여한다. 사회적 실재감(social presence)은 매체를 통한 상호작용에서 상대방을 실제 존재로 지각하는 정도를 가리키며[27], Lombard와 Ditton[28]은 이를 현존감(presence) 개념으로 발전시켜 매체 환경에서의 몰입도와 실재감 지각이 사회적 상호작용의 질과 신뢰 형성에 미치는 영향을 분석하였다.

국내 실증 연구는 의인화에서 신뢰로 이어지는 매개 경로를 정교하게 검증해 왔다. 양현란과 지성구[9]는 쇼핑 챗봇에서 의인화가 사회적 실재감과 라포를 거쳐 신뢰와 만족·재이용의도로 이어지는 경로를 구조방정식으로 확인하였으며, 최상묵과 최도영[29]은 이커머스 챗봇 환경에서 챗봇 신뢰가 쇼핑몰·판매자·브랜드 신뢰로 전이되며 이 전이 과정을 사회적 실재감이 조절함을 보고하였다. 유희승[7]은 AI 휴먼 챗봇 이용자의 스피치 역량과 인지된 상호작용 품질이 신뢰도 형성을 매개함을 확인하였다.

교육 분야 아바타 에이전트의 효과에 관한 최근 연구는 구현 효과(embodiment effect)의 실증적 근거를 축적해 왔다. Zhang과 Pan[8]은 2014년부터 2024년까지의 교육 분야

대화형 에이전트 문헌에 대한 범위 검토를 수행하여 시각·청각 구현이 학습 참여와 사회적 실재감을 촉진함을 종합하였다. Mayer와 DaPra[30]는 컴퓨터 기반 학습 환경에서 아바타 구현이 학습자의 이해도와 학습 효과를 유의하게 높임을 실험적으로 입증하였다. Pataranutaporn 외[31]는 실제 인물의 음성·외모를 모사한 디지털 트윈이 학습자의 정서적 지원에 활용될 수 있음을 보고하여, 아바타가 단순 의인화를 넘어 정체성 이전(identity transfer)의 매체로 기능할 수 있음을 시사하였다.

2-4 신뢰 지향 설계 접근의 필요성

이상의 검토를 종합하면, AI 아바타 챗봇의 신뢰성은 ① 챗봇 정체성, ② 답변 품질, ③ 대화 자연스러움, ④ 운영 정책 투명성이라는 복수의 축에 걸친 다층적 신호 설계에 의해 결정된다. 그러나 선행 연구는 이러한 축을 개별적으로 검증해왔으며, 이들을 단일 시스템 안에 통합 설계하고 그 통합 설계의 사용자 인지 효과를 실증한 사례는 부족하다.

본 연구는 이러한 공백을 메우기 위해 신뢰를 설계 단계에서 의도적으로 구조화하는 신뢰 지향 접근을 제안한다. 본 접근은 다층의 신뢰 신호를 단일 시스템에 통합 배치한 뒤 사용자 인지를 통해 그 설계의 도달성을 실증한다. 본 연구는 이 접근을 대학 신입생의 전공 선택·진로 상담 챗봇에 적용하여 4개 계층 18개 컴포넌트로 구성된 설계 체계를 제안한다.

본 분석 틀의 핵심은 이론적 근거·시스템 구현·설문 측정을 동일 식별자 아래 1:1로 연결한다는 점에 있다. 한 신뢰 컴포넌트가 어떤 이론에 근거하고 시스템의 어떤 위치에 구현되며 사용자에게 어떻게 측정되는지를 일관되게 추적할 수 있으며, 시스템 구현 가능성과 사용자 인식 도달성을 같은 분석 단위에서 함께 점검할 수 있다. 본 분석 틀의 구체적 구성은 3장에서 상세히 기술한다.

III. 연구 방법

3-1 시스템 설계

1) 4계층 아키텍처 개요

본 시스템은 그림 1과 같이 사용자 단말, 엣지 프록시, 대화 코어, 그리고 학교 서버의 네 개 계층으로 구성된다(표 1).

사용자가 화면에서 입력한 메시지는 T1(브라우저)에서 T2(엣지 프록시)를 거쳐 T3(자체 GPU 서버)으로 전달되며, 여기서 RAG 검색과 LLM 추론이 수행된다. 인증·로그·설문 데이터는 T4(학교 서버)에서 별도로 보관된다. 네 계층은 서로 다른 호스트와 프로그램 환경에서 운영되며, 한 계층의 내부 구현이 변경되어도 다른 계층의 동작에 영향을 주지 않도록 HTTP(S) 인터페이스만으로 통신한다.

이러한 계층 분리는 세 가지 설계 의도를 반영한다. 첫째, 학생의 식별 정보와 발화 데이터를 학교 서버(T4)에서만 다루어 외부 클라우드 사업자에게 사용자 데이터가 누적되지 않도록 한다. 둘째, LLM 추론과 임베딩 연산을 자체 GPU 서버(T3)에서 수행하여 외부 추론 API에 대한 의존을 구조적으로 통제하고 데이터 주권을 학내에 유지한다. 셋째, 인증·추론·영속화의 책임을 서로 다른 계층에 분산하여 한 계층의 장애가 전체 시스템의 데이터 무결성으로 전파되지 않도록 한다.

System Architecture Overview

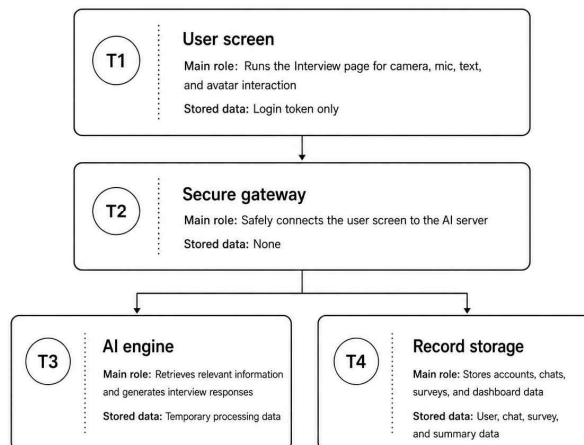


그림 1. 4계층 시스템 아키텍처

Fig. 1. Four-tier system architecture

표 1. 4계층 시스템 아키텍처 구성

Table 1. Four-tier system architecture configuration

Stage	Plain name	Main role	Stored data
T1	User screen	Runs the interview page for camera, mic, text input, and avatar interaction	Login token only
T2	Secure gateway	Safely connects the user screen to the AI server	None
T3	AI engine	Retrieves relevant information and generates interview responses	Temporary processing data
T4	Record storage	Stores accounts, chats, surveys, and dashboard data	User, chat, survey, and summary data

2) 검색 증강 생성(RAG) 파이프라인

T3의 RAG 파이프라인은 본 시스템 답변의 사실성을 담보하는 핵심 장치이다. 본 연구는 OpenAI·Google 등 외부 LLM API를 사용하지 않고 자체 GPU 호스트에서 검색과 생성을 모두 수행하여, 학생의 질문·발화가 외부 클라우드로 전달되지 않도록 설계하였다. 도메인 코퍼스는 C대학 미래융합대학 11개 전공의 공식 소개 영상을 전사하여 질문·답변 쌍으로 가공한 후, 학생 의견을 반영하여 학과 사무실 연락처를 보강한 131개 청크로 구성된다. 각 청크는 식별자·전공 분류·질문·답변·키워드를 포함한 JSON 객체로 저장된다.

검색 절차는 다음과 같다. 사용자가 질문을 입력하면 임베딩 모델[25]을 통해 질문이 1024차원 벡터로 변환되고, 인텍스의 청크 벡터들과 코사인 유사도가 계산된다. 유사도 임계값(0.25) 이상인 청크 중 상위 5개가 시스템 프롬프트의 참고 자료 절에 삽입되어 LLM(Gemma4, Ollama 런타임)에 전달된다. 검색 결과가 부족한 경우에는 모델이 임의로 추론하지 않고 “교수님께 직접 여쭙보세요” 등의 안내문을 출력하도록 강제되며, 이는 환각을 억제하는 핵심 안전장치이다 [4],[24].

생성된 응답은 8단계 후처리를 거친다. 화면 표시용 텍스트와 음성 발화용 텍스트를 분리한 뒤, 이모지·장식 부호 제거, 연락처 정돈, 마무리 문장 정규화, 영문 약어와 전공명의 음성 발음 치환(예: “AI” → “에이아이”), 한국어 수사 정규화, 복합 명사의 음절 분리 순으로 처리된다. 이러한 후처리는 동일한 사실 정보가 음성과 화면 두 채널에 각각 적합한 형태로 전달되도록 보장한다.

3) 멀티모달 인터페이스 통합 (FTF·STS·TTT)

본 시스템의 사용자 인터페이스는 그림 2와 같이 좌측의 영상 상호작용 영역과 우측의 대화 영역으로 분리되며, 화면 하단에서 HeyGen LiveKit Stream, WebSocket, REST API, 그리고 STT 파이프라인이 연결되어 영상·음성·텍스트 입력이 단일 상담 흐름 안에서 통합 처리된다. 좌측 영상 영

역은 AI 아바타 패널과 사용자 카메라 패널로 구성되어 음성·영상 입출력을 담당하고, 우측 대화 영역은 메시지 스트림, 컨트롤 카드, 입력 도크로 구성되어 텍스트 입출력을 담당한다.

본 시스템은 단일 백엔드에서 세 가지 대화 모드를 통합 지원한다. FTF(Face-to-Face)는 음성으로 묻고 아바타 영상으로 답을 받는 대면 모드, STS(Speech-to-Speech)는 음성으로 주고받는 음성 전용 모드, TTT(Text-to-Text)는 키보드로 주고받는 텍스트 전용 모드이다. 사용자는 한 세션 안에서 자유롭게 모드를 전환할 수 있다. 세 모드 모두 동일한 RAG·LLM·후처리 경로를 거치므로 어떤 모드를 선택하더라도 답변의 사실 정확성은 동일하게 유지된다[8]. 다만 음성 발화용 텍스트와 화면 표시용 텍스트는 별도 필드로 분리되어, 화면에서는 전화번호·이메일이 클릭 가능한 카드로 표시되는 반면 음성에서는 “학과 홈페이지” 같은 자연어로 발화된다.

영상 상호작용 모듈의 세부 구성은 그림 3에 제시하였다. AI 아바타 패널은 실제 교수자의 외모·음성을 모사한 HeyGen 아바타가 음성으로 응답하는 영역이며, 사용자 카메라 패널은 사용자의 카메라 입력을 받는 영역이다. 사용자는 카메라 토글, 마이크 토글, 세션 종료 버튼을 통해 상담 진행 상태를 직접 제어할 수 있고, 연결 상태 표시기가 실시간 통신 상태를 안내한다. 음성 출력과 음성 입력은 양방향 순환 구조로 작동하여 단순한 영상 표시 장치가 아닌 실시간 대화 인터페이스로 기능한다.

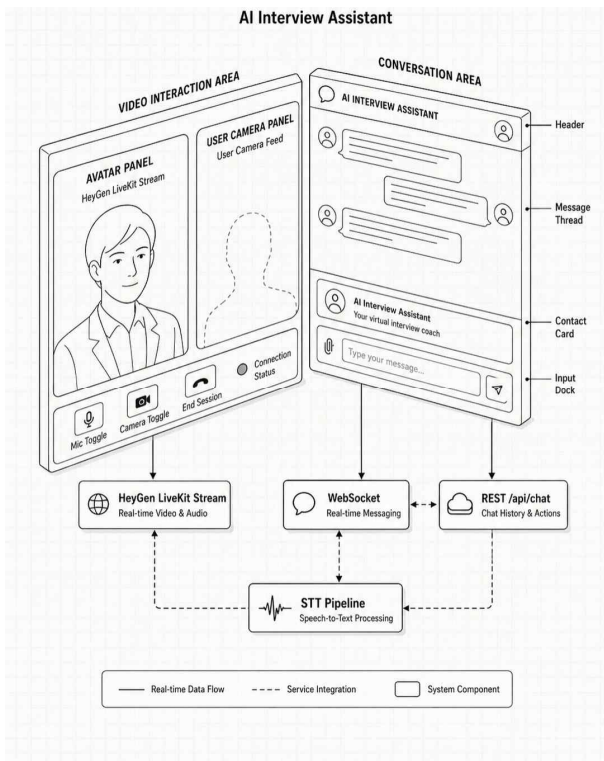


그림 2. AI 면접 어시스턴트의 전체 사용자 인터페이스 구조
Fig. 2. Overall user interface architecture of the AI interview assistant

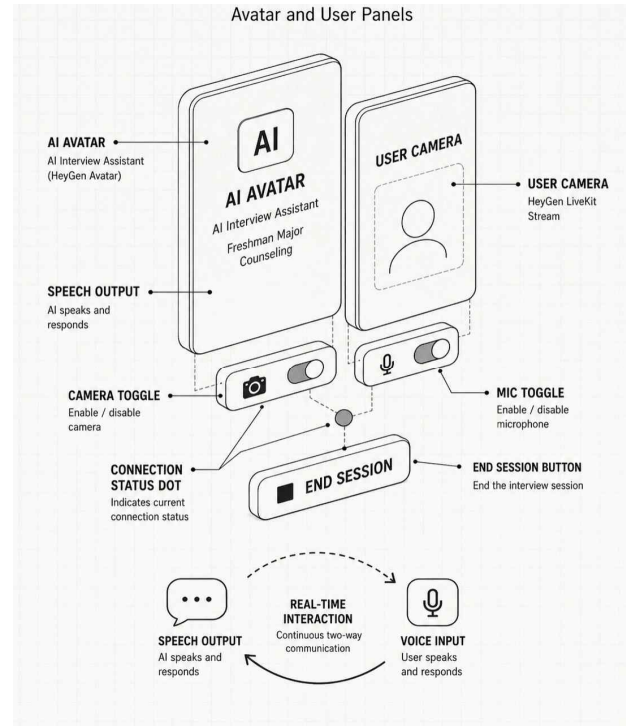


그림 3. 영상 상호작용 모듈의 구성 (아바타 및 사용자 카메라 패널)
Fig. 3. Video interaction module with avatar and user camera panels

채팅 상호작용 모듈의 세부 구성은 그림 4에 제시하였다. 채팅 패널은 상단 헤더, 중앙 메시지 스레드, 하단 입력 도크의 세 영역으로 나뉜다. 헤더에는 설문 버튼·다크모드 전환·사용자 정보·로그아웃 기능이 배치되고, 메시지 스레드에는 AI 응답 말풍선, 사용자 발화 말풍선, 그리고 컨택 정보를 시각화한 정보 카드가 표시된다. 입력 도크에는 마이크 버튼, 텍스트 입력창, 전송 버튼이 포함되어 음성 입력과 텍스트 입력을 모두 지원한다.

학과 사무실 정보를 질의한 경우, 시스템은 사용자 메시지에 명시된 전공명을 길이 내림차순으로 매칭하여 해당 학과의 전화번호·홈페이지·학과장 이메일을 추출하고 응답 옆에 클릭 가능한 컨택 카드(그림 4의 Info Card)로 표시한다. 매칭 우선순위를 검색 점수가 아닌 사용자 메시지에 둔 이유는, 검색 점수만으로는 "미술치료 사무실"과 같은 질의가 의미적으로 가까운 다른 학과(예: 심리학)로 오분류될 수 있기 때문이다. 컨택 카드는 모바일과 데스크톱 환경 모두에서 동일한 클릭 동작(전화 걸기·메일 작성·외부 페이지 열기)을 제공하여 채널 간 경험을 일관되게 유지한다.

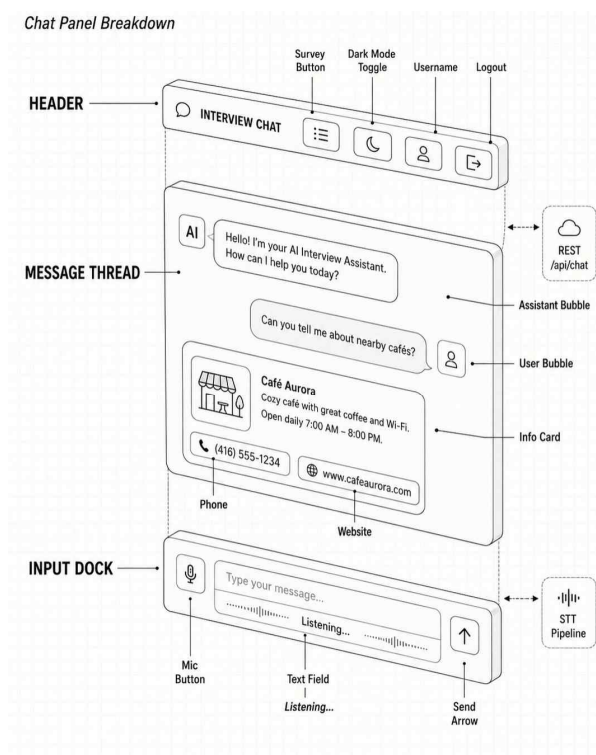


그림 4. 채팅 상호작용 모듈의 구성(메시지 스레드 및 입력 도크)
Fig. 4. Chat interaction module and message input structure

3-2 4계층 18 신뢰 컴포넌트

본 연구는 2장에서 정리한 신뢰 이론을 단일 챗봇 시스템에 체계적으로 담기 위해, 4개 계층에 걸친 18개 컴포넌트로 구성된 신뢰 지향 설계 체계를 제안한다(표 2·3·4·5). 4계층

구분은 신뢰 신호의 성격(정체성·내용·상호작용·정책)과 사용자에게 보이는 가시성을 함께 반영한다.

L1 정체성(F1-F3, 3개)은 “이 봇이 누구의 어떤 도구인가”를 사용자에게 알리는 신호로, Mayer 외[12]의 ABI 모형에서 진실성(integrity) 차원과 Hovland와 Weiss[32]가 제안한 출처 신뢰성(source credibility)에 대응한다. L2 콘텐츠 품질(F4-F7, 4개)은 봇이 무엇을 어떻게 말하는가의 영역으로, ABI의 능력(ability) 차원과 김민성·구철모[16]가 보고한 인지적·정서적 신뢰의 분화에 대응한다. L3 인터랙션 미시(F8-F12, 5개)는 평소에는 잘 드러나지 않다가 작동에 문제가 생길 때 비로소 인지되는 신호로, CASA·MASA 패러다임[17],[18]과 사용자 통제권[33]을 다룬다. L4 메타 신호(F13-F18, 6개)는 단발 상호작용이 아닌 누적 이용 경험과 운영 정책에서 형성되는 학습된 신뢰[14] 및 초기 신뢰[15]를 겨냥한다.

표 2. 4계층 × 18 신뢰 컴포넌트 매핑과 인지율(L1)

Table 2. Mapping of the 18 trust components across four layers, with awareness rates (L1)

Layer	ID	Trust signal	Functional description	Theoretical basis
L1	F1	Full digital twin	Avatar aligned with the real instructor's appearance, voice, speech style, and answer content	Doney & Cannon [34]; Pataranutaporn et al. [31]
L1	F2	Institutional identity display	Operating institution consistently shown in header, meta-tags, and greeting	Hovland & Weiss [32]; Bhattacharjee [35]
L1	F3	AI identity disclosure	“AI assistant” explicitly stated in greeting and on the avatar nameplate	Liao & Sundar [26]

표 3. 4계층 × 18 신뢰 컴포넌트 매핑과 인지율(L2)

Table 3. Mapping of the 18 trust components across four layers, with awareness rates (L2)

Layer	ID	Trust signal	Functional description	Theoretical basis
L2	F4	RAG-based factuality	Responses grounded in retrieved domain chunks to suppress hallucination	Ji et al. [22]
L2	F5	Explicit limitation acknowledgment	Blocks speculative inference and emits a guidance message when retrieval falls below threshold	Lee & See [20]; Jacovi et al. [21]
L2	F6	Warm honorific tone	Affective warmth via Korean polite-honorific style and student-centered phrasing	Pelau et al. [36]; Lombard & Xu [18]
L2	F7	Output format consistency	Consistent response format with decorative symbols removed	Ji et al. [22]

본 체계의 18개 컴포넌트는 각각 시스템의 기능 단위와 1:1로 대응되고, 평가 설문 18문항(F1~F18)과도 동일한 식별자로 연결되어 있다. 따라서 한 컴포넌트의 이론적 근거 → 시스템 구현 → 이용자 인식 측정이 같은 ID 아래 일관되

표 4. 4계층 × 18 신뢰 컴포넌트 매핑과 인지율(L3)

Table 4. Mapping of the 18 trust components across four layers, with awareness rates (L3)

Layer	ID	Trust signal	Functional description	Theoretical basis
L3	F8	Response latency management	Per-device silence-threshold tuning for natural turn-taking	Skantze [37]; Cowan et al. [38]
L3	F9	Multi-layer echo guard	Prevents the bot's own speech from re-entering the user microphone	Skantze [37]
L3	F10	ESC speech interrupt	ESC key immediately interrupts bot speech, granting user control	Buçinca et al. [33]; Shneiderman [39]
L3	F11	Avatar lip-sync & gesture	Visual-auditory embodiment that enhances social presence	Bickmore & Picard [40]; Mayer & DaPra [30]
L3	F12	Smooth mode transition	Free transition among FTF, STS, and TTT modalities	Walther [41]

표 5. 4계층 × 18 신뢰 컴포넌트 매핑과 인지율(L4)

Table 5. Mapping of the 18 trust components across four layers, with awareness rates (L4)

Layer	ID	Trust signal	Functional description	Theoretical basis
L4	F13	Privacy consent UI	Consent screen explicitly stating collected items, purpose, and retention period	Acquisti et al. [42]; Liao & Sundar [26]
L4	F14	Guest browsing	No-signup entry point for low-commitment exploration	McKnight et al. [15]
L4	F15	Korean ordinal greeting	Visit count expressed naturally as Korean ordinal numbers	Sundar & Marathe [43]
L4	F16	Visit-count tracking	Greeting that acknowledges cumulative visits for returning users	Horton & Wohl [44]; Lombard & Xu [18]
L4	F17	TTS pronunciation normalization	Pronunciation normalization for English acronyms, institution names, and major-name compounds	(operational signal)
L4	F18	KakaoTalk in-app external redirect	Auto-redirect from KakaoTalk in-app browser to external browser	Shneiderman [39]

게 추적된다.

각 컴포넌트의 'Yes %'는 본 분석 표본(N = 50) 중 해당 컴포넌트에 응답 가능한 사용자가 'Yes'로 답한 비율이다. 일부 컴포넌트(F9·F11·F16·F17·F18)는 음성·영상 모드 사용 경험, 재방문 여부, 카카오톡 인앱 진입 등의 사전 조건이 충족된 사용자만 응답하도록 설계되었다. 4계층 평균은 컴포넌트별 Yes 응답률의 산술평균으로 정의된다. 각 컴포넌트의 시스템 구현 코드 위치 및 설문 문항 원문은 부록 A에 수록한다.

각 계층의 핵심 차별점은 다음과 같다. L1에서 가장 강한 차별 요소는 F1 폴 디지털 트윈이다. 시스템은 해당 교수의 외모·음성·말투를 외부 아바타 서비스의 Studio Avatar(직접 녹화), 음성 클로닝, 본인 Q&A 기반 RAG 청크, 시스템 프롬프트의 본인 말투 규칙, 디지털 트윈 명시의 다섯 차원에서 정합시켰다. L2에서는 RAG 기반 사실성과 정직한 한계 인정이 인지적 차원을, 해요체의 따뜻한 말투가 정서적 차원을 담당한다. L3는 평소에 가시화되지 않는 견고성 신호로, 사용자가 상호작용을 자연스러운 대화로 지각하도록 한다. L4는 운영 정책 차원의 신호로, 개인정보 처리 투명성·신규 사용자 진입로·언어적 배려·재방문 인자·매체 환경 일관성을 포괄한다.

3-3 설문 설계 및 분석 절차

본 연구는 4계층 18개 컴포넌트에 대한 이용자 인식을 측정하기 위해 탐색적 횡단 설문을 수행하였다. 카카오톡 채널을 통해 배포된 시스템 링크로 본 챗봇을 이용한 응답자를 자율 참여 형식으로 모집하였다.

설문지는 인구통계 5문항(학년·성별·MBTI·1전공·2전공), 신뢰 컴포넌트 18문항(F1~F18), 전반 신뢰 1문항(Q24), 자유응답 2문항(Q25 신뢰가 갔던 순간·Q26 신뢰가 떨어진 순간)의 총 26문항으로 구성된다. 응답 형식은 인구통계를 제외한 모든 컴포넌트 문항에 대해 Yes/No 이분형으로 통일하였다. 이분형 채택의 이유는 두 가지로, 세션 직후 즉시 노출되는 설문에서 응답 부담을 최소화하여 완료율을 확보하고, 각 컴포넌트가 “이번 세션에서 해당 신뢰 신호를 인지하였는가”라는 단일 이진 판단으로 자연스럽게 조작화될 수 있다고 보았기 때문이다. 다만 이분형은 신뢰의 정도 차이를 포착하지 못한다는 측정학적 한계가 있으며, 본 한계는 5장에서 다시 논의한다.

18개 컴포넌트 중 5개(F9·F11·F16·F17·F18)는 응답자가 전제조건(음성 또는 영상 모드 사용 경험, 재방문 여부, 카카오톡 인앱 진입 여부)을 충족하지 못한 경우 클라이언트 측에서 자동으로 비활성 처리되어 NULL로 저장된다. 설문 모달은 사용자 발화 누적 3턴 이상에 도달한 세션이 종료되는 시점에 자동으로 표시되거나 사용자가 임의로 호출할 수 있다.

표본 정제 절차는 다음과 같다. 자율 참여로 수집된 총 69명의 응답 중, 응답 소요 시간이 60초 미만인 부주의 응답(flag_too_fast=1) 17명을 1차로 제외하였다. 이후 잔여 응답 시간 분포에서 임계값에 가장 근접한 두 응답을 추가로 제

외하여 최종 분석 표본을 N = 50으로 확정하였다. NULL 응답은 분모에서 제외되어 응답 가능 여부와 무관하게 균일한 분모가 적용될 때 발생하는 천장 효과를 회피하였다.

분석은 본 연구의 탐색적 성격에 부합하도록 기술 통계에 한정한다. 구체적으로는 (i) 18개 컴포넌트 각각의 Yes 응답 비율, (ii) 4개 계층(L1~L4)별 평균 점수를 산출한다. 변수 간 인과 관계 추론과 다변량 분석은 본 연구의 범위에서 제외한다. 자유응답(Q25, Q26) 텍스트는 두 명의 연구자가 네 개의 사용자 경험 카테고리(① 기술적 한계, ② 정보 깊이의 부족, ③ 분위기 불일치, ④ 발화 주체에 대한 통합 인식 부족)로 분류하여 5장 논의에서 결과 해석의 근거로 활용하며, 측정 도구 자체의 개념 매핑상 주의 사항은 5장 한계에서 별도로 다룬다. 본 분류는 학내 동료 검토를 거쳐 보정되었다.

IV. 연구 결과

4-1 표본 기술

분석 표본은 본 챗봇을 이용한 C대학 신입생 50명이다. 신입생은 전공 선택과 진로 상담의 의사결정 부담이 가장 큰 집단으로, 본 연구가 설정한 전공 선택·진로 상담 도메인의 핵심 사용자에게 해당한다. 표본 정제 절차는 3-3절에 따랐으며, 본 연구는 학년·성별·전공별 비교 분석을 목적으로 하지 않으므로 인구통계 분포는 별도로 보고하지 않는다.

4-2 컴포넌트별 인지도

18개 신뢰 컴포넌트와 전반 신뢰(Q24)의 Yes 응답 비율을 그림 5와 표 2에 정리하였다. 컴포넌트별 인지도율은 50%대부터 100%까지 큰 편차를 보였다. 천장 효과(95% 이상)가 관찰된 컴포넌트는 F18 카카오톡 인앱 진입자에게만 응답이 활성화되는 조건부 컴포넌트이며, F13은 본 시스템의 회원가입 동의 UI와 카카오톡 OAuth 동의 화면이 함께 노출된 결과로 측정되었다. 두 인터페이스 중 어느 쪽이 응답자의 인지에 더 크게 기여했는지는 본 설문 도구로는 분리되지 않으므로, 후속 측정에서 두 동의 화면을 분리해 묻는 문항 보강이 필요하다.

가장 낮은 인지율을 보인 컴포넌트는 F2 기관 신뢰 표시(52.0%), F1 디지털 트윈(56.0%), F11 아바타 입동가·제스처(67.3%), F9 에코 가드(69.4%) 순이었다. 이들은 각각 정체성 신호의 두 항목과 인터랙션 미시 신호의 두 항목으로, 일부 신호가 사용자 인지로 충분히 도달하지 못한 양상을 드러낸다. 본 격차의 원인은 5장에서 네 가지 사용자 경험 요인으로 분석한다.

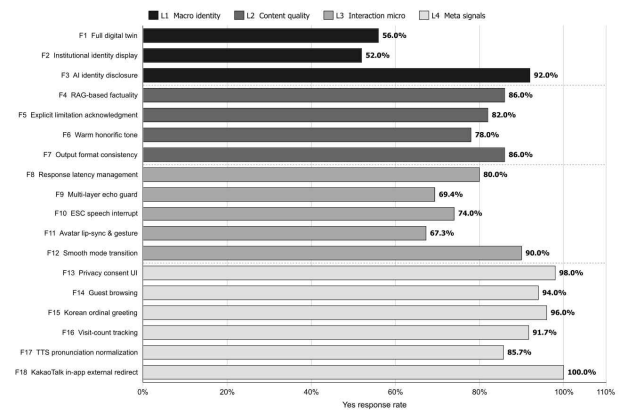


그림 5. 18개 신뢰 컴포넌트별 Yes 응답률

Fig. 5. Awareness rate of the 18 trust components

4-3 4계층 비교 - 인식 수준 격차

4개 계층의 평균 인지율은 다음과 같이 산출되었다. L1 정체성 평균 66.7%(범위 52~92%), L2 콘텐츠 품질 평균 83.0%(범위 78~86%), L3 인터랙션 미시 평균 76.1%(범위 67~90%), L4 메타 신호 평균 94.2%(범위 86~100%). 즉 본 연구가 중점적으로 설계한 정체성 계층이 운영 정책 차원의 메타 신호 계층에 비해 27.5퍼센트 포인트 낮은 인지율을 보이는 인식 수준 격차(perception gap)이 관찰되었다(그림 6).

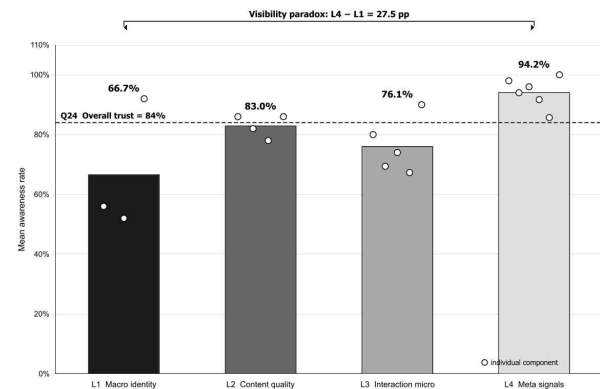


그림 6. 4계층 평균 인지율과 인식 수준 격차

Fig. 6. Layer-mean awareness rate and the perception gap

L1 계층 내부에서도 인지율이 균질하지 않았다. AI 정체를 명시적 라벨(이름판)로 노출한 F3은 92.0%의 높은 인지율을 보인 반면, 붓의 정체성을 분산된 다중 단서(아바타 영상·음성·말투·RAG 답변·시스템 프롬프트)로 구현한 F1과 헤더·메타 정보·인사말로 분산 노출한 F2은 각각 56.0%, 52.0%로 절반 수준에 머물렀다. 이는 명시적 단서는 잘 인지되지만 분산된 다중 단서는 통합 인지로 이어지지 않는 패턴을 시사하며, Pataranutaporn 외[31]가 제시한 디지털 트윈 효과가 작동하기 위해서는 단서들이 사용자 측에서 통합 인지될 만큼 명시적이거나 임계값을 넘어야 함을 함의한다.

4-4 신뢰 인식의 세부 패턴

L2 콘텐츠 품질 계층 내부에서는 인지적 차원과 정서적 차원의 인지율 비대칭이 관찰되었다. 정확성과 정직성을 측정하는 인지적 차원의 컴포넌트인 F4 RAG 사실성(86.0%), F5 한계 인정(82.0%), F7 형식 일관성(86.0%)의 평균은 84.7%였다. 반면 따뜻함과 친근함이라는 정서적 차원을 측정하는 F6 따뜻한 말투는 78.0%로, 인지적 차원과 6.7퍼센트 포인트의 차이를 보였다. 본 비대칭은 김민성과 구철모[16]가 ChatGPT 사용 맥락에서 보고한 인지적 신뢰와 정서적 신뢰의 분화와 부분적으로 정합하며, 본 시스템이 LLM 기반 답변에서 정확성과 정직성은 비교적 안정적으로 사용자에게 도달시킨 반면 정서적 차원은 동일한 강도로 도달시키지 못하였음을 보여준다.

전반 신뢰(Q24)는 84.0%로 측정되었으며, 이는 L2 콘텐츠 품질 계층의 평균 인지율(83.0%)과 분포적으로 거의 일치한다. 다른 계층의 평균과의 차이는 L1과 17.3퍼센트 포인트, L3과 7.9퍼센트 포인트, L4와 10.2퍼센트 포인트로 모두 L2와의 정합도가 가장 높았다. 본 분포 정합은 인과 관계를 의미하지 않으나, 사용자가 붓 전체에 대한 신뢰를 평가할 때 콘텐츠 품질 차원에 가장 큰 비중을 두는 경향을 가설적으로 시사하며, 후속 연구의 인과 분석에서 검증해야 할 가설로 제안된다.

V. 논의 및 결론

5-1 점수 격차의 가능한 설명

4장의 결과에서 관찰된 인지율 격차(인식 수준 격차, L2 인지/정서 비대칭, 일부 컴포넌트의 낮은 점수)는 단일 원인이 아닌 복수의 사용자 경험 요인으로 설명될 수 있다. 본 절은 자유응답(Q25·Q26) 텍스트의 사후 분류에 근거하여 네 가지 사용자 경험 요인(① 기술적 한계, ② 정보 깊이 부족, ③ 분위기 불일치, ④ 발화 주체에 대한 통합 인식 부족)을 제시한다. 이와 별개로 측정 도구 자체의 개념 매핑상 주의 사항은 5-3절(연구의 한계)에서 함께 다룬다.

1) 기술적 한계

시스템 구현이 의도된 신호를 충분히 송출하지 못한 경우이다. 자유 응답에서 다수 응답자가 영상·음성 채널의 비동기를 직접 지적하였다.

“말 소리와 입모양의 시차 발생.”

“발음이 약간 끊겨지는 부분이 있었고... 처음 이용하는 학생의 경우 교수님의 모습이 나오는 부분에서 AI 상담인지 실제 상담인지 구분이 어려워 초기 혼란을 느낄 가능성이 있다.”

“말 중간에 오류” · “중간에 멈춤이나 언어 선택이 부자연스러웠다.”

음성 인식 측면에서도 “음성으로 질문할 때 말이 아직 안 끝났는데 다 했다고 인지할 때”와 같은 사용자 발화 차단이 보고되었다. 이러한 기술적 한계는 F1(56.0%)·F6(78.0%)·F8(80.0%)·F9(69.4%)·F11(67.3%)의 인지율 저하와 직결되며, 본 시스템이 외부 상용 아바타 서비스에 의존한 구조적 제약과 무관하지 않다.

2) 정보 깊이의 부족

RAG 코퍼스의 도메인 정보가 사용자 기대 수준에 미치지 못한 경우이다. 한 응답자는 환각과 오분류 사례를 직접 보고하였다.

“설득의 이해” 등 존재하지 않는 교과목을 수강하라고 추천.”

“심리학과 경영학의 융합을 물어봤는데 갑자기 미디어커뮤니케이션 전공을 추천.”

회피성·표층적 응답에 대한 지적(“회피성답변으로 대답하여 신뢰가 떨어졌다”, “두루뭉술한 대답만 들었습니다”, “전공 조합에 대해 물었을 때 답변이 애매하다고 느꼈다”)도 반복적으로 등장하였다. 또한 “자격정보, 졸업 후 진로, 실습 및 비교과 활동, 교수진의 전문분야 등에 대한 내용이 보다 구체적으로 반영되면 활용도가 높아질 것 같다”는 보장 요청은, 본 시스템이 학사 정보의 표층적 깊이는 다루었으나 진로·실습·인지 자원 정보로 확장되지 못했음을 시사한다. 본 카테고리에는 F2(52.0%), F4 일부, F5(82.0%, 회피성으로 오해), F6 일부의 인지율에 영향을 미친 것으로 해석된다.

3) 분위기 불일치

외양·음성은 모방했으나 인격적 톤은 충분히 이전되지 못한 경우이다. 해당 교수를 평소에 알고 있는 한 응답자는 본 시스템의 차분한 응답 톤이 실제 인물의 텐션과 불일치함을 직접 지적하였다.

“제가 평소에 느끼는 교수님은 정말 친절하고 텐션이 높은데요. 교수님을 알고 있는 사람이 이번 붓을 경험하면 어색함을 느낄 것 같습니다!”

“그래도 아직은 완벽히 교수님 같지 않음”이라는 보다 일반적인 지적도 같은 맥락이다. 본 카테고리는 F1(56.0%)과 F6(78.0%)의 인지율 저하에 직접적으로 작용하며, Pataranutaporn 외[31]가 제시한 디지털 트윈 효과의 경계 조건(외양·음성 모방이 인격적 정체성 이전을 자동으로 보장하지 않는다는 점)을 실증적으로 드러낸다.

4) 발화 주체에 대한 통합 인식 부족

붓의 발화 주체를 가리키는 신호가 여러 차원에 흩뿌려져 있어 사용자가 한 덩어리로 통합 인식하지 못한 경우이다. L1 계층 내부에서 명시적 라벨로 노출된 F3(AI 정체 명시, 92.0%)과 분산된 다중 단서로 구현된 F1·F2(56.0%, 52.0%) 사이에 36~40퍼센트 포인트의 격차가 관찰되었으며, 이는 노

출 방식의 차이가 인지율에 결정적 영향을 미쳤음을 시사한다.

디지털 트윈은 외양·음성·말투·RAG 답변·시스템 프롬프트의 다섯 차원에 분산 구현되어 있고, 기관 신원은 헤더·메타·인사말에 분산되어 있다. 단서들의 합이 단서를 통합하는 사용자의 인지 작업을 자동으로 보완해 주지는 않았다. 본 카테고리 F1과 F2의 인지율 저하를 가장 강하게 설명한다.

요컨대 본 결과의 인식 수준 격차와 L2 비대칭은 위 네 가지 요인이 개별적으로, 또는 복합적으로 작용한 결과로 해석된다. 특히 가장 낮은 인지율을 보인 F2(52.0%)과 F1(56.0%)은 ②·③·④가 중첩되어 영향을 미친 사례로 보인다. 다만 위 분류는 자유응답에 기반한 사후적 해석으로, 통제된 실험을 통한 각 요인의 효과 검증은 후속 연구의 과제이다. 본 다각적 해석은 신뢰 지향 설계의 시스템 구현뿐 아니라 만들어진 신호가 사용자에게 어떻게 도달했는지를 함께 검증해야 함을 시사한다.

5-2 결과 요약과 함의

본 연구는 신뢰 이론을 종합한 4계층 18 컴포넌트 신뢰 지향 설계 체계를 멀티모달 아바타 챗봇에 매핑·구현하고, 50명의 이용자 설문을 통해 그 도달성을 평가하였다. RQ1에 대한 답으로, 단일 백엔드 위에서 FTF·STS·TTT 세 모달리티를 통합 운영하는 4계층 아키텍처를 통해 모달리티 간 사실성 편차를 완화하면서 인터페이스 신호와 콘텐츠 품질을 동시에 담보하는 통합 설계가 실제로 구현 가능함을 보였다. 사용자의 전반 신뢰는 84.0%, 메타 신호 계층의 평균 인지율은 94.2%에 이르렀다. RQ2에 대한 답으로는, 18개 신호가 모두 동일한 수준으로 사용자에게 전달되는 것은 아니었다. 정체성 계층의 인지율(66.7%)이 메타 신호 계층(94.2%)에 비해 27.5퍼센트 포인트 낮은 인식 수준 격차가 관찰되었으며, 콘텐츠 품질 계층 내부에서도 인지적 차원(84.7%)에 비해 정서적 차원(78.0%)이 약하게 전달되는 비대칭이 확인되었다. 사용자의 전반 신뢰(84.0%)는 콘텐츠 품질 계층의 평균(83.0%)과 분포적으로 가장 정합하였다.

이론적 함의로, 신뢰 형성에서 시스템 구현뿐 아니라 만들어진 신호가 사용자에게 어떻게 도달했는지를 함께 검증해야 한다는 점이 가장 중요하다. 즉 잘 만들었다고 자동으로 신뢰가 형성되는 것이 아니라, 사용자 인식에 도달한 신호를 별도로 점검해야 한다. 본 연구의 4계층 18 컴포넌트 체계는 이를 동일 식별자 아래 추적할 수 있는 분석 도구로 활용될 수 있으며, 5-1절의 네 카테고리 분류는 격차의 원인 진단을 위한 휴리스틱으로 기능한다.

실무적 함의로, 디지털 트윈 신호의 전달력을 강화하기 위해 분산 배치된 단서를 사용자가 통합 인식하도록 유도하는 설계 개선이 요청된다. 첫 발화에서 디지털 트윈의 다섯 차원이 통합되어 있음을 명시적으로 안내하거나, 사용자가 붓과 해당 교수의 정체성 정합을 확인할 수 있는 가시적 표시를 추가하는 방안이 가능하다. 콘텐츠 품질 내 정서적 차원의 전달력 강화를 위해서는 해요체와 따뜻한 말투의 일관성을 시스

템 프롬프트 차원에서 강화하고 RAG 청크의 정서적 표현 풍부도를 검토할 필요가 있다.

5-3 연구의 한계

본 연구의 한계는 표본의 외적 타당도, 측정·분석 도구의 정밀성, 시스템 재현성의 세 측면에서 후속 연구의 보완을 요구한다. 첫째, 표본의 외적 타당도가 제한적이다. N = 50의 탐색적 표본 규모와 자율 참여 표집 방식은 본 연구의 일반화 가능성에 본질적 제약을 가한다. 본 연구의 응답자는 카카오톡 채널을 통해 배포된 시스템 링크에 자율적으로 접속한 학생들로, 시스템에 관심을 보인 학생만이 링크를 통해 자발적으로 접속하므로, 호의적인 사용자가 표본에 우선 포함되는 편향이 구조적으로 내재되어 있다. 모집 채널 자체가 학과 카카오톡 채널이라는 점에서 학과 정체성에 대한 사전 인지가 상대적으로 높은 응답자가 표본에 과대 대표되었을 가능성도 있다. 따라서 본 연구에서 보고된 4계층 평균 인지율(L1 66.7%, L2 83.0%, L3 76.1%, L4 94.2%)과 전반 신뢰(84.0%)는 일반 신입생 모집단의 인식 수준에 비해 상향 편향되어 측정되었을 가능성을 배제할 수 없다. 또한 단일 시점의 횡단 설계로 누적 사용 경험에 따른 학습된 신뢰의 동태적 변화를 추적하지 못하였다. 후속 연구에서는 학사 시스템 연동을 통한 무작위 표집 또는 강제 노출 설계로 자기선택 편향을 통제하고, 종단 설계로 학습된 신뢰의 형성 과정을 추적할 필요가 있다.

둘째, 측정 도구와 분석 방법의 정밀성이 제한적이다. Yes/No 이분형 응답은 정도 차이를 포착하지 못하여 일부 자유응답에서 관찰된 경계 인식(“발음이 뭉개지는 부분이 있으나 크게 문제가 될 정도는 아니었습니다”)을 반영하기 어렵고, 단일 문항의 전반 신뢰 측정(Q24)은 다차원 신뢰 측정의 깊이를 확보하지 못하였다. 또한 일부 컴포넌트는 신뢰의 인접 구성개념을 측정한다는 점에 유의해야 한다. F2가 측정된 “공식 도구로 느꼈는가”는 기관 신뢰와 정확히 일치하지 않으며, F1이 측정된 “본인과 대화하는 듯한 느낌”은 동일시(identification)에 가깝고 붓임을 알면서도 답이 쏠리면 신뢰하는 적정 신뢰(calibrated trust)[20]와는 구별된다. F10은 실제 사용 경험과 기능 존재 추측을 문항 표현상 분리하지 못하였다. 분석 방법 또한 기술 통계에 한정되어 변수 간 인과 관계 추론과 다변량 분석을 시도하지 않았다.

셋째, 시스템의 재현 가능성이 외부 의존성에 의해 제약된다. 본 시스템은 외부 상용 아바타 서비스(HeyGen Interactive Avatar)에 의존하며, 본 API의 운영 종료 일정이 사전 공지된 상황에서 운영 연속성과 재현 가능성에 구조적 위험이 내포되어 있다. 본 의존성은 5-1절 ① 기술적 한계에서 보고된 립싱크 비동기 등 다수의 사용자 경험 문제와도 직접 연결되어 있어, 신뢰 신호의 전달성에 본 외부 시스템이 미치는 영향을 본 연구의 설계 변수만으로 통제할 수 없었다.

5-4 후속 연구 방향

본 연구의 한계는 후속 연구의 출발점이 된다. 첫째, 표본 확장과 인과 분석으로의 전환이다. 표본을 N=100 이상으로 확장하여 4계층 체계의 구조 타당도를 확인적 요인분석으로 검증하고, 4계층 점수를 독립변수로 한 로지스틱 회귀를 통해 전반 신뢰의 선행 변수를 식별할 수 있다. 본 18개의 컴포넌트 외의 추가 컴포넌트를 확충하고 분석 패러다임을 기술 통계에서 인과 분석으로 전환하여 신뢰 형성의 인과 구조를 본격 검증할 수 있다. 둘째, 중단 설계의 도입이다. 동일 이용자의 반복 사용 경험에 따른 학습된 신뢰의 동태적 형성을 중단(longitudinal) 설계로 추적할 수 있다. 셋째, 측정 도구의 정교화이다. 다항 척도와 reverse-coded 문항을 도입하고 인지적·정서적 신뢰의 분화[16]를 본 체계에 통합 측정함으로써 신뢰의 다차원 구조를 정밀하게 검증할 수 있다. 넷째, 도메인 및 시스템 일반화이다. 본 4계층 체계를 의료·정신 건강 상담 등 다른 도메인에 적용하여 일반화 가능성을 검증하는 동시에, 외부 상용 아바타 의존을 완화하는 차세대 WebRTC 기반 자체 호스팅 스택으로 마이그레이션하여 운영 연속성과 재현 가능성을 확보할 수 있다.

본 연구가 제안한 신뢰 지향 4계층 체계와 멀티모달 시스템 아키텍처는 전공 선택·진로 상담 도메인의 한 사례에 머물지 않고, 도메인 사실성과 인터페이스 일관성을 동시에 요구하는 모든 대화형 AI 시스템의 신뢰 지향 설계 출발점으로 검토 가능성을 지닌다. 후속 연구는 이러한 일반화의 경험적 근거를 축적하는 데 기여할 수 있을 것이다.

감사의 글

This research includes results partly supported by the “Gyeonggi Regional Innovation System & Education Project (Gyeonggi RISE Project)”, supported by the Ministry of Education and Gyeonggi Province (grant No. 20250093)

부록. 설문 문항(총 26문항)

본 연구의 자체 설문지는 4부 26문항으로 구성된다. 본문의 신뢰 컴포넌트 식별자 F1~F18은 설문 문항 Q06~Q23과 1:1 대응된다.

1. 신뢰 컴포넌트 인지(Q06~Q23 / F1~F18)

각 문항은 “이번 세션에서 다음 신호를 인지하셨습니까?”의 형식으로 Yes/No 이분형으로 응답한다.

표 6. 신뢰 컴포넌트 인지 설문 문항(Q06~Q23 / F1~F18)

Table 6. Trust-component awareness survey items (Q06~Q23 / F1~F18)

Survey ID	In-text ID	Question (English reference translation)
Q06	F1	Did the bot feel like talking with the actual instructor?
Q07	F2	Did the bot feel like an official tool of C University?
Q08	F3	Was it clearly conveyed that the bot is an AI assistant rather than a human?
Q09	F4	Did the bot's responses feel grounded in actual departmental information?
Q10	F5	Did the bot honestly admit limitations when it did not know?
Q11	F6	Did the bot's tone feel warm and friendly?
Q12	F7	Did the bot's response format appear consistent and professional?
Q13	F8	Did the bot's response speed maintain a natural conversational flow?
Q14 †	F9	Did the bot avoid mistakenly recognizing its own speech as user input?
Q15	F10	Did you feel that you could interrupt the bot when it spoke for a long time?
Q16 †	F11	Were the avatar's lip-sync and gestures natural?
Q17	F12	Was the transition among voice, video, and text modes smooth?
Q18	F13	At sign-up, were the data-collection items and purposes clearly explained?
Q19	F14	Was an option to browse the bot without signing up provided?
Q20	F15	Did you receive a Korean ordinal greeting such as “Welcome on your first visit”?
Q21 †	F16	On revisit, did you receive a greeting that acknowledged your cumulative visits?
Q22 †	F17	Were English acronyms and major names pronounced naturally?
Q23 †	F18	When entering through KakaoTalk, did the system auto-redirect to an external browser?

† 조건부 문항: 다음 사전 조건이 충족된 응답자만 활성화. F9~F17은 음성 모드 사용자, F11은 영상(FTF) 모드 사용자, F16은 재방문(2회 이상), F18은 카카오톡 인앱 진입자.

2. 전반 신뢰(Q24)

Q24. 종합적으로 본 봇의 답변을 신뢰할 수 있다고 느끼셨나요? (Yes/No)

3. 자유응답(Q25, Q26)

Q25. (선택 입력, 2,000자 이내) 가장 신뢰가 갔던 순간이 나 기능은 무엇이었나요?

Q26. (선택 입력, 2,000자 이내) 반대로 가장 신뢰가 떨어졌거나 어색했던 순간은 언제였나요?

4. 응답 처리 절차

설문 모달은 사용자 발화 누적 3턴 이상에 도달한 세션이 종료되는 시점에 자동으로 표시되거나 사용자가 임의로 호출할 수 있다. 응답 소요 시간이 60초 미만이면 부주의 응답으로 표지(flag_too_fast)되어 분석에서 제외된다. 조건부 문항의 미충족 응답은 NULL로 저장되어 분모에서 제외된다.

참고문헌

- [1] M. Dawood, "Assessing the Effectiveness of Chatbots in Providing Personalized Academic Advising and Support to Higher Education Students: A Narrative Literature Review," *Studies in Technology Enhanced Learning*, Vol. 4, No. 1, October 2024. <https://doi.org/10.21428/8c225f6e.7140f8f4>
- [2] G. Bilquise, S. Ibrahim, and S. M. Salhieh, "Investigating Student Acceptance of an Academic Advising Chatbot in Higher Education Institutions," *Education and Information Technologies*, Vol. 29, No. 5, pp. 6357-6382, 2024. <https://doi.org/10.1007/s10639-023-12076-x>
- [3] M. A. Kuhail, N. Alturki, S. Alramlawi, and K. Alhejori, "Interacting with Educational Chatbots: A Systematic Review," *Education and Information Technologies*, Vol. 28, No. 1, pp. 973-1018, 2023. <https://doi.org/10.1007/s10639-022-11177-3>
- [4] J. Swacha and M. Gracel, "Retrieval-Augmented Generation (RAG) Chatbots for Education: A Survey of Applications," *Applied Sciences*, Vol. 15, No. 8, 4234, April 2025. <https://doi.org/10.3390/app15084234>
- [5] Z. Li, Z. Wang, W. Wang, K. Hung, H. Xie, and F. L. Wang, "Retrieval-Augmented Generation for Educational Application: A Systematic Survey," *Computers and Education: Artificial Intelligence*, Vol. 8, 100417, 2025. <https://doi.org/10.1016/j.caeai.2025.100417>
- [6] H. Yusuf, A. Money, and D. Daylamani-Zad, "Pedagogical AI Conversational Agents in Higher Education: A Conceptual Framework and Survey of the State of the Art," *Educational Technology Research and Development*, Vol. 73, No. 2, pp. 815-874, 2025. <https://doi.org/10.1007/s11423-025-10447-4>
- [7] H. S. Yu, "The Process of Trust Formation in AI Human Chatbots: The Serial Multiple Mediation Effect of Opinion Leaders' and General Users' Speech Competence and Perceived Interaction Quality," *Knowledge & Liberal Arts*, Vol. 17, pp. 1083-1126, March 2025. <https://doi.org/10.54698/kl.2025.17.1083>
- [8] Y. Zhang and W. Pan, "A Scoping Review of Embodied Conversational Agents in Education: Trends and Innovations from 2014 To 2024," *Interactive Learning Environments*, Vol. 33, No. 7, pp. 4566-4587, 2025. <https://doi.org/10.1080/10494820.2025.2468972>
- [9] H. L. Yang and S. G. Ji, "Antecedents and Consequences of Customer Rapport for Shopping Chatbots," *Journal of Distribution and Logistics*, Vol. 8, No. 4, pp. 5-21, December 2021. <https://doi.org/10.22321/jdl2021080401>
- [10] H. S. Jeong and Y. I. Kim, "The Effect of Chatbot Quality on Chatbot Trust and Brand Trust," *Journal of the Korean Society of Costume*, Vol. 69, No. 3, pp. 1-14, 2019. <https://doi.org/10.7233/jksc.2019.69.3.001>
- [11] S. H. Lim, K. K. Lee, and J. Y. Lee, "An Empirical Study on the Utilization of Generative AI in University Students' Learning: Focusing on Generative AI Chatbot Service Quality, Trusts, Satisfaction, and Intention to Continue Use," *Journal of Product Research*, Vol. 42, No. 4, pp. 1-7, 2024. <https://doi.org/10.36345/kacst.2024.42.4.001>
- [12] R. C. Mayer, J. H. Davis, and F. D. Schoorman, "An Integrative Model of Organizational Trust," *The Academy of Management Review*, Vol. 20, No. 3, pp. 709-734, July 1995. <https://doi.org/10.5465/amr.1995.9508080335>
- [13] Y. Lee, H. Kim, and M. Park, "The Effects of Perceived Quality of Fashion Chatbot's Product Recommendation Service on Perceived Usefulness, Trust and Consumer Response," *Journal of the Korean Society of Clothing and Textiles*, Vol. 46, No. 1, pp. 80-98, February 2022. <https://doi.org/10.5850/JKSC.2022.46.1.80>
- [14] K. A. Hoff and M. Bashir, "Trust in Automation: Integrating Empirical Evidence on Factors That Influence Trust," *Human Factors: The Journal of the Human Factors and Ergonomics Society*, Vol. 57, No. 3, pp. 407-434, 2015. <https://doi.org/10.1177/0018720814547570>
- [15] D. H. McKnight, V. Choudhury, and C. Kacmar, "Developing and Validating Trust Measures for E-Commerce: An Integrative Typology," *Information Systems Research*, Vol. 13, No. 3, pp. 334-359, September 2002. <https://doi.org/10.1287/isre.13.3.334.81>
- [16] M. Kim and C. Koo, "A Study of the Behavioral Intention on Conversational ChatGPT for Tourism Information Search Service: Focusing on the Role of Cognitive and

- Affective Trust,” *Information Systems Review*, Vol. 26, No. 1, pp. 119-149, February 2024. <https://doi.org/10.14329/isr.2024.26.1.119>
- [17] C. Nass, J. Steuer, and E. R. Tauber, “Computers Are Social Actors,” in *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, Boston: MA, pp. 72-78, April 1994. <https://doi.org/10.1145/191666.191703>
- [18] M. Lombard and K. Xu, “Social Responses to Media Technologies in the 21st Century: The Media Are Social Actors Paradigm,” *Human-Machine Communication*, Vol. 2, pp. 29-55, 2021. <https://doi.org/10.30658/hmc.2.2>
- [19] S.-Y. Park and S. Lee, “A Study on the Effect of Intimacy Between Conversational Agents and Users on Reliability: Focused on Self Exposure, Small Talk and Anthropomorphism,” *Journal of Korea Design Forum*, Vol. 24, No. 3, pp. 55-62, 2019. <https://doi.org/10.21326/ksdt.2019.24.3.005>
- [20] J. D. Lee and K. A. See, “Trust in Automation: Designing for Appropriate Reliance,” *Human Factors*, Vol. 46, No. 1, pp. 50-80, 2004. <https://doi.org/10.1518/hfes.46.1.50.30392>
- [21] A. Jacovi, A. Marasović, T. Miller, and Y. Goldberg, “Formalizing Trust in Artificial Intelligence: Prerequisites, Causes and Goals of Human Trust in AI,” in *Proceedings of The 2021 ACM Conference on Fairness, Accountability, and Transparency*, Virtual Event, Canada, pp. 624-635, March 2021. <https://doi.org/10.1145/3442188.3445923>
- [22] Z. Ji, N. Lee, R. Frieske, T. Yu, D. Su, Y. Xu, ... and P. Fung, “Survey of Hallucination in Natural Language Generation,” *ACM Computing Surveys*, Vol. 55, No. 12, 248, March 2023. <https://doi.org/10.1145/3571730>
- [23] P. Lewis, E. Perez, A. Piktus, F. Petroni, V. Karpukhin, N. Goyal, ... and D. Kiela, “Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks,” in *Proceedings of the 34th International Conference on Neural Information Processing Systems*, Vancouver, Canada, pp. 9459-9474, December 2020. <https://doi.org/10.48550/arXiv.2005.11401>
- [24] W. Zhang and J. Zhang, “Hallucination Mitigation for Retrieval-Augmented Large Language Models: A Review,” *Mathematics*, Vol. 13, No. 5, 856, March 2025. <https://doi.org/10.3390/math13050856>
- [25] J. Chen, S. Xiao, P. Zhang, K. Luo, D. Lian, and Z. Liu, “M3-Embedding: Multi-Linguality, Multi-Functionality, Multi-Granularity Text Embeddings through Self-Knowledge Distillation,” in *Findings of the Association for Computational Linguistics: ACL 2024*, Bangkok, Thailand, pp. 2318-2335, August 2024. <https://doi.org/10.18653/v1/2024.findings-acl.137>
- [26] Q. V. Liao and S. S. Sundar, “Designing For Responsible Trust in AI Systems: A Communication Perspective,” in *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*, Seoul, Korea, pp. 1257-1268, June 2022. <https://doi.org/10.1145/3531146.3533182>
- [27] J. Short, E. Williams, and B. Christie, *The Social Psychology of Telecommunications*, London, UK: Wiley, 1976.
- [28] M. Lombard and T. Ditton, “At the Heart of It All: The Concept of Presence,” *Journal of Computer-Mediated Communication*, Vol. 3, No. 2, September 1997. <https://doi.org/10.1111/j.1083-6101.1997.tb00072.x>
- [29] S.-M. Choi and D. Choi, “The Effect of Customer Experience on Trust Transfer in E-Commerce Chatbot Environment: Focusing on the Moderating Effect of Social Presence,” *The Journal of the Korea Contents Association*, Vol. 22, No. 7, pp. 136-148, July 2022. <https://doi.org/10.5392/JKCA.2022.22.07.136>
- [30] R. E. Mayer and C. S. DaPra, “An Embodiment Effect in Computer-Based Learning with Animated Pedagogical Agents,” *Journal of Experimental Psychology: Applied*, Vol. 18, No. 3, pp. 239-252, 2012. <https://doi.org/10.1037/a0028616>
- [31] P. Pataranutaporn, V. Danry, J. Leong, P. Punpongsanon, D. Novy, P. Maes, and M. Sra, “AI-Generated Characters for Supporting Personalized Learning and Well-Being,” *Nature Machine Intelligence*, Vol. 3, No. 12, pp. 1013-1022, December 2021. <https://doi.org/10.1038/s42256-021-00417-9>
- [32] C. I. Hovland and W. Weiss, “The Influence of Source Credibility on Communication Effectiveness,” *Public Opinion Quarterly*, Vol. 15, No. 4, pp. 635-650, 1951. <https://doi.org/10.1086/266350>
- [33] Z. Buçinca, M. B. Malaya, and K. Z. Gajos, “To Trust or to Think: Cognitive Forcing Functions Can Reduce Overreliance on AI in AI-Assisted Decision-Making,” *Proceedings of the ACM on Human-Computer Interaction*, Vol. 5, No. CSCW1, 188, pp. 1-21, April 2021. <https://doi.org/10.1145/3449287>
- [34] P. M. Doney and J. P. Cannon, “An Examination of the Nature of Trust in Buyer-Seller Relationships,” *Journal of Marketing*, Vol. 61, No. 2, pp. 35-51, April 1997. <https://doi.org/10.1177/002224299706100203>
- [35] A. Bhattacharjee, “Individual Trust in Online Firms: Scale Development and Initial Test,” *Journal of Management Information Systems*, Vol. 19, No. 1, pp. 211-241, 2002. <https://doi.org/10.1080/07421222.2002.11045715>

- [36] C. Pelau, D.-C. Dabija, and I. Ene, "What Makes An AI Device Human-Like? The Role of Interaction Quality, Empathy and Perceived Psychological Anthropomorphic Characteristics in the Acceptance of Artificial Intelligence in the Service Industry," *Computers in Human Behavior*, Vol. 122, 106855, September 2021. <https://doi.org/10.1016/j.chb.2021.106855>
- [37] G. Skantze, "Turn-Taking In Conversational Systems and Human-Robot Interaction: A Review," *Computer Speech & Language*, Vol. 67, 101178, May 2021. <https://doi.org/10.1016/j.csl.2020.101178>
- [38] B. R. Cowan, N. Pantidi, D. Coyle, K. Morrissey, P. Clarke, S. Al-Shehri, ... and N. Bandeir, "What Can I Help You With?": Infrequent Users' Experiences of Intelligent Personal Assistants," in *Proceedings of the 19th International Conference on Human-Computer Interaction with Mobile Devices and Services*, Vienna, Austria, 43, pp. 1-12, September 2017. <https://doi.org/10.1145/3098279.3098539>
- [39] B. Shneiderman, "Human-Centered Artificial Intelligence: Reliable, Safe & Trustworthy," *International Journal of Human-Computer Interaction*, Vol. 36, No. 6, pp. 495-504, 2020. <https://doi.org/10.1080/10447318.2020.1741118>
- [40] T. W. Bickmore and R. W. Picard, "Establishing and Maintaining Long-Term Human-Computer Relationships," *ACM Transactions on Computer-Human Interaction*, Vol. 12, No. 2, pp. 293-327, June 2005. <https://doi.org/10.1145/1067860.1067867>
- [41] J. B. Walther, "Computer-Mediated Communication: Impersonal, Interpersonal, and Hyperpersonal Interaction," *Communication Research*, Vol. 23, No. 1, pp. 3-43, February 1996. <https://doi.org/10.1177/009365096023001001>
- [42] A. Acquisti, L. Brandimarte, and G. Loewenstein, "Privacy and Human Behavior in the Age of Information," *Science*, Vol. 347, No. 6221, pp. 509-514, January 2015. <https://doi.org/10.1126/science.aaa1465>
- [43] S. S. Sundar and S. S. Marathe, "Personalization versus Customization: The Importance of Agency, Privacy, and Power Usage," *Human Communication Research*, Vol. 36, No. 3, pp. 298-322, July 2010. <https://doi.org/10.1111/j.1468-2958.2010.01377.x>
- [44] D. Horton and R. R. Wohl, "Mass Communication and Para-Social Interaction: Observations on Intimacy at a Distance," *Psychiatry*, Vol. 19, No. 3, pp. 215-229, 1956. <https://doi.org/10.1080/00332747.1956.11023049>

성봉주(Bong Ju Sung)



2025년 : 차의과학대학교
데이터경영학과 (경영학사)

2025년~현 재: 차의과학대학교 일반대학원

AI헬스케어융합학과 석사과정

※ 관심분야 : AI 헬스케어(AI Healthcare), 신뢰 지향 AI 시스템(Trustworthy AI Systems), 멀티모달 아바타 챗봇(Multimodal Avatar Chatbots), 검색 증강 생성(Retrieval-Augmented Generation, RAG), 디지털 헬스케어 콘텐츠(Digital Healthcare Content)

이규성(Gyu Seong Lee)



2025년 : 차의과학대학교
데이터경영학과 (경영학사)

2025년~현 재: 차의과학대학교 일반대학원

AI헬스케어융합학과 석사과정

※ 관심분야 : 멀티모달·멀티태스크 학습(Multimodal/Multi-task Learning), 생성형 인공지능(Generative AI), 약동학·집단약동학(PK/PopPK), 의료 인공지능(Medical AI), 의료 데이터 분석(Medical Data Analysis) 등

김정환(Junghwan Kim)



2008년 : 연세대학교
커뮤니케이션대학원
(영상학석사)

2014년 : 연세대학교
커뮤니케이션대학원
(영상학박사)

2017년~현 재: 차의과학대학교 미디어커뮤니케이션학 전공
주임교수

2024년~현 재: 차의과학대학교 일반대학원

AI헬스케어융합학과 교수

※ 관심분야 : AI 매개 커뮤니케이션(AI-Mediated Communication), 비주얼 커뮤니케이션(Visual Communication), 미디어 심리학(Media Psychology), AI 미디어 콘텐츠 제작(AI Media Content Production)



박대근(Dae Keun Park)

1996년 : 한국과학기술원 (산업경영학
석사)

2008년 : 한국과학기술원
(경영공학박사)

2000년 ~ 2004년: 한국신용정보(NICE)

2005년 ~ 2009년: 삼정KPMG

2010년 ~ 2014년: 액센츄어

2016년 ~ 현 재: 차의과학대학교 경영학 전공 주임교수

2024년 ~ 현 재: 차의과학대학교 일반대학원

AI헬스케어융합학과 교수

※ 관심분야 : 재무금융, 시계열 예측, 생성형AI, 헬스케어융합 등