

AI 헬스 윤리 리터러시 개념 정립 및 항목 도출을 위한 델파이 연구

이 하 나¹ · 엄 주 희^{2*}

¹이화여자대학교 연령통합고령사회연구소 연구교수

²건국대학교 융합인재학부 교수

A Delphi Study on the Conceptualization and Item Generation of AI Health Ethics Literacy

Hannah Lee¹ · Juhee Eom^{2*}

¹Research Professor, Ewha Institute for Age Integration Research, Seoul 03760, Korea

²Professor, Department of Law, KU GLOCAL Innovation, Chungju 27478, Korea

[요 약]

디지털 기기와 AI 기술의 발전은 의료 혁신을 이끌며 예방적 건강 관리 패러다임을 제시하고 있다. 그러나 AI 헬스케어 기술의 확산은 알고리즘 편향, 개인정보 보호 등 윤리적 문제를 동반한다. 기존 연구가 의료 제공자와 정책 중심의 윤리 논의에 집중한 반면, 본 연구는 사용자 중심의 AI 헬스 윤리 리터러시(AI Health Ethics Literacy)를 개념화하고 이를 평가할 수 있는 초기 척도 항목을 탐색적으로 도출하는 것을 목표로 하였다. 이를 위해 델파이 기법을 활용해 AI·의료·윤리 전문가 15명을 대상으로 세 차례 조사를 시행하였으며, 개인정보 보호, 알고리즘 공정성, 투명성 등 12개 핵심 요소를 포함하는 척도를 도출하였다. 본 연구에서 도출한 척도는 사용자가 AI 기반 의료기술 이용 시 직면할 수 있는 주요 윤리적 이슈들을 구체화하여, AI 헬스 윤리 리터러시의 개념적 틀을 정립하는 데 기여한다. 또한, 개인정보 보호 수준, 알고리즘의 공정성과 투명성에 대한 사용자 인식을 체계적으로 측정할 수 있는 항목들을 제시함으로써, 향후 AI 헬스케어 기술의 윤리적 수용성을 평가하고 사용자 신뢰 형성에 실질적으로 기여할 수 있는 기초 자료를 제공한다는 점에서 의의를 갖는다.

[Abstract]

The advancement of digital devices and artificial intelligence (AI) has driven healthcare innovation, introducing a paradigm of preventive health management. However, the expansion of AI healthcare technologies raises ethical concerns, including algorithmic bias and data privacy. While previous research focused on provider- and policy-centered ethical discussions, this study conceptualizes user-centered AI Health Ethics Literacy and explores the key components for developing an assessment scale. Using the Delphi method, three rounds of surveys were conducted with 15 experts in AI, healthcare, and ethics, leading to the identification of 12 key components, such as privacy protection, algorithmic fairness, and transparency. The proposed assessment scale concretizes critical ethical issues that users may face when engaging with AI-based healthcare technologies, thereby contributing to the establishment of a conceptual framework for AI Health Ethics Literacy. Moreover, by systematically presenting items to evaluate users' awareness of privacy protection, algorithmic fairness, and transparency, this study provides foundational resources for assessing the ethical acceptability of AI healthcare technologies and substantially contributes to fostering user trust.

색인어 : AI 헬스 윤리 리터러시, AI 윤리, 디지털 헬스, 척도 개발, 델파이 조사

Keyword : AI Health Ethics Literacy, AI Ethics, Digital Health, Scale Development, Delphi Method

<http://dx.doi.org/10.9728/dcs.2025.26.4.1091>



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 February 2025; **Revised** 12 March 2025

Accepted 25 March 2025

***Corresponding Author, Juhee Eom**

Tel: +82-43-840-3348

E-mail: juheelight@kku.ac.kr

I. 서 론

디지털 기기와 인공지능(AI) 기술의 발전은 의료 전반에 혁신적인 변화를 가져왔다. AI 기반 기술은 질병 예측·진단·치료뿐만 아니라 일상적 영역으로 확장되면서, 예방적 건강 관리라는 새로운 패러다임을 제시하고 있다[1]. 예를 들어, 웨어러블 디바이스는 사용자의 생체 신호를 실시간으로 모니터링하여 건강 상태를 관리하고, 개인 맞춤형 헬스케어 서비스를 제공한다[2]. 또한, AI 기반 영상 분석 기술은 방대한 의료 데이터를 신속하게 분석하여 질병을 조기에 진단하고, 개인 맞춤형 치료 전략을 수립하는 데 기여한다[3]. 이러한 기술은 기존 의료의 한계를 넘어 지속적인 건강 증진 개입과 건강 관리 방식에 근본적인 변화를 일으키고 있다.

그러나 이러한 긍정적 측면에도 불구하고, AI 기반 헬스케어 기술의 발전은 알고리즘 편향, 개인정보 유출, 오진과 같은 다양한 윤리적 문제를 동반한다[4],[5]. 특정 인구 집단에 불리한 결과를 초래하는 알고리즘 편향[6], 의료 정보의 부적절한 보호로 인한 개인정보 유출[7] 등의 문제는 AI 헬스케어 기술이 신뢰를 얻는 데 있어 중요한 장애 요소가 된다. AI 기술이 환자의 건강과 직접적으로 연관된 민감한 정보를 다루고, 진단 및 건강 관리 의사 결정 과정에 깊이 개입하는 만큼, 사용자 관점의 윤리적 고려가 무엇보다 중요하다.

이러한 맥락에서 본 연구는 AI 헬스 윤리 리터러시(AI Health Ethics Literacy) 개념을 정립하고, 사용자가 AI 헬스케어 기술을 이용하는 과정에서 직면하는 윤리적 이슈를 평가할 수 있는 기준을 마련하는 데 초점을 둔다. AI 헬스케어 기술의 보편화는 윤리적 고려를 바탕으로 사용자 신뢰를 구축하는 것을 요구하지만, 아직 사용자 관점에서 AI 헬스 윤리에 대한 인식을 체계적으로 평가한 연구는 부족한 실정이다. 기존 연구들은 주로 의료 서비스 제공자 관점에서 윤리적 문제를 논의하거나 정책적 차원의 규제와 가이드라인 수립에 초점을 맞추어 왔다[8],[9]. 이러한 연구들은 AI 기술의 윤리적 문제를 기술적·법적 관점에서 논의하는 데 중요한 기여를 했으나, 사용자가 직면하는 AI 헬스케어 윤리 문제를 체계적으로 다루지 못했다는 제한점이 있다.

AI 헬스케어 기술이 일상적인 건강관리 수단으로 자리 잡으면서, 사용자가 이러한 기술의 윤리적 요소를 이해하고 비판적으로 평가하는 것이 더욱 중요해지고 있다. 특히, 개인정보 보호, 알고리즘의 공정성과 투명성, 책임성과 같은 윤리적 문제는 사용자들이 반드시 이해하고 평가해야 할 핵심 요소들이다[7],[10]. AI 헬스케어 기술의 신뢰성과 윤리적 사용을 촉진하기 위해서는 사용자의 윤리적 인식을 체계적으로 평가할 수 있는 개념적 틀이 필요하다. 이와 관련하여 기존의 디지털 헬스 리터러시(digital health literacy) 개념은 사용자 역량을 평가하는 유용한 기초를 제공한다.

디지털 헬스 리터러시는 주로 건강 정보에 대한 접근·이해·활용 능력을 중심으로 개인의 건강 관리 역량을 측정하고 평가하는 데 사용되어 왔다[11]~[13]. 그러나 AI 헬스케어 기

술이 도입됨에 따라 새롭게 부각된 윤리적 과제들은 기존 디지털 헬스 리터러시 척도의 한계를 드러내고 있다. 기존 척도는 개인정보 보호와 같은 일부 윤리적 요소를 다루고 있지만[13], AI 의사결정 과정의 투명성, 알고리즘 편향의 인지, AI 기반 건강 관리 기술의 신뢰성 평가 등의 문제를 충분히 포괄하지 못한다. 이러한 한계는 AI 헬스케어 기술의 특수성을 반영하여 사용자 관점에서 윤리적 문제를 체계적으로 평가할 수 있는 새로운 척도의 필요성을 시사한다. 즉, 기존 척도가 일반적인 디지털 환경에서의 건강 정보 탐색 및 활용이라는 개인의 정보 처리 역량에 초점을 맞추었다면, 본 연구는 알고리즘 편향성, AI 의사결정 과정의 투명성, 책임성 등 AI 헬스케어 기술과 관련된 윤리적 쟁점들을 통합적으로 평가할 수 있는 척도를 제안한다는 점에서 차별성을 지닌다.

본 연구는 이러한 필요성을 충족하기 위해 AI 헬스 윤리 리터러시(AI Health Ethics Literacy)의 개념을 구체화하고, 윤리적 인식을 평가할 수 있는 초기 항목을 도출하는 것을 목표로 한다. AI 헬스케어 기술 도입이 국내에서도 빠르게 진행되고 있지만, 관련 윤리에 대한 학문적 논의는 여전히 부족한 실정이다. 특히, 한국형 AI 헬스 윤리 리터러시 척도의 개발은 국내 환경에 적합한 맞춤형 윤리 전략을 수립하기 위한 필수적인 선행 과제이다. AI 헬스케어 기술에 대한 사용자 인식은 사회문화적 맥락의 영향을 받기 때문에[14],[15], 본 연구는 서구권에서 개발된 AI 윤리 척도를 참고하되, 한국적 특수성을 반영한 초기 항목을 도출하고자 한다. 이를 위해 델파이 기법을 활용하여 전문가 의견을 수렴하고, 공정성·투명성·개인정보 보호 등 주요 윤리적 기준이 한국 문화적 맥락에 적합하도록 반영할 계획이다. 이러한 척도는 단순히 소비자의 인식을 파악하는 데 그치지 않고, AI 헬스케어 기술의 설계와 운영 과정에서 실질적인 변화를 이끌어내며, 궁극적으로 사용자의 건강 증진에 기여할 것으로 기대한다.

II. 문헌고찰

2-1 AI 기반 헬스케어 기술과 디지털 헬스 윤리

현대 의료 환경에서 AI 기반 헬스케어 기술은 방대한 의료 데이터를 수집·분석하여 개인 맞춤형 건강 관리와 치료 계획을 지원하는 혁신적 도구로 자리 잡았다. 이러한 기술 발전은 진단 정확도 향상, 예방적 건강 관리, 맞춤 치료 등의 긍정적 효과를 제공하는 동시에, 개인정보 보호, 데이터 투명성, 공정성 및 편향, 책임 소재 등 복합적인 윤리적 문제를 수반한다[16],[17]. 따라서 기술의 신뢰성과 사회적 수용성을 확보하기 위해서는, AI 헬스케어 기술이 내포한 윤리적 쟁점을 체계적으로 이해하고 관리하는 것이 필수적이다.

특히, 개인정보 보호와 데이터 투명성은 AI 기반 시스템이

환자의 민감한 건강 데이터를 수집·분석하는 과정에서 중요한 고려 대상이다. Trocin et al.[16]은 데이터 처리 과정에서 프라이버시 침해 위험을 지적하며, 투명성과 보안을 보장할 수 있는 체계적 접근의 필요성을 강조하였다. 또한, Möllmann et al.[18]은 AI 시스템이 ‘블랙박스’로 작동함에 따라 의료 의사결정 과정에서 신뢰 저하가 발생할 수 있음을 경고하였고, Maeckelberghe et al.[19]은 설명 가능한 AI(Explainable AI, XAI)를 도입하여 사용자가 의사결정의 근거를 명확히 이해할 수 있도록 할 것을 제안하였다.

또한, 공정성과 편향의 문제도 중요한 윤리적 과제로 대두된다. AI가 학습 데이터에 내재한 편향을 반영할 경우, 특정 인구 집단에 불공정한 결과를 초래할 위험이 있다[16],[20]. Möllmann et al.[18]은 데이터 수집 및 처리 과정에서 공정성을 유지하기 위한 엄격한 관리와 모니터링의 필요성을 강조하였으며, Maeckelberghe et al.[19]은 국제적 표준 및 정책 도입을 통해 다양한 문화적·사회적 맥락에서도 공정하게 작동할 수 있는 AI 시스템 구축의 필요성을 제기하였다.

책임 소재 문제 역시 AI 기술의 핵심 윤리적 이슈로 꼽힌다. 의료 사고 발생 시 설계자, 의료진, 시스템 운영자 등 누구에게 책임을 물어야 하는지 명확히 규정되지 않으면, 기술에 대한 신뢰와 사회적 수용성이 저하될 수 있다[16],[18]. 이와 같이, AI 기반 헬스케어 기술이 야기하는 다양한 윤리적 쟁점을 종합적으로 고찰할 때, 단순한 정보 소비를 넘어 공정성·투명성·책임성·개인정보 보호 등 윤리적 요소를 체계적으로 평가하고 교육할 수 있는 통합적 접근의 필요성이 더욱 부각된다. 이러한 통합적 프레임워크는 AI 헬스케어 기술의 신뢰성과 사회적 수용성을 높이는 데 기여할 수 있는 견고한 기반을 마련할 것이다.

2-2 기존 디지털 헬스 리터러시 척도의 한계와 AI 헬스 윤리 리터러시로의 확장

디지털 헬스 리터러시는 개인이 디지털 환경에서 의료 정보를 탐색·평가·통합·활용할 수 있는 능력을 의미하며, 디지털 기술의 발전과 함께 의료 정보 접근 및 관리의 핵심 개념으로 자리 잡았다. Norman과 Skinner[21]는 전자 건강 리터러시를 정보 문해력, 건강 문해력, 컴퓨터 문해력 등 여섯 가지 핵심 요소로 구성된 다차원적 개념으로 정의하고, 이를 평가하기 위한 척도로 eHEALS(eHealth Literacy Scale)을 개발하였다. 이 척도는 사용자의 정보 탐색 능력, 신뢰성 평가, 디지털 기술 활용 능력 등 정보 소비 중심의 역량을 측정하는 데 중요한 역할을 해왔으나, 디지털 기술이 Health 1.0(정보 소비 중심)에서 Health 2.0(상호작용 및 자가 건강 관리 중심)으로 진화함에 따라 기존 eHEALS는 새로운 상호작용 및 참여 요구를 충분히 반영하지 못하는 한계를 드러내게 되었다. 이에 따라 Van der Vaart와 Drossaert[13]는 기존 척도의 한계를 보완하기 위해 DHLI(Digital Health Literacy Instrument)를 개발하였으며, DHLI는 탐색 기술, 신뢰성 평

가, 개인정보 보호, 사용자 생성 콘텐츠 관리와 같은 기술적 요소를 포함하여 디지털 환경에서 요구되는 복합적인 역량을 평가할 수 있는 확장된 틀을 제공하였다.

그러나 eHEALS와 DHLI는 디지털 헬스 리터러시의 측정 도구로서 중요한 역할을 해왔음에도 불구하고, AI 의료기술 환경에서 새롭게 부각되는 윤리적 문제들을 충분히 반영하지 못하는 한계를 지니고 있다. Siau와 Wang[20]은 AI 기술 확산이 알고리즘 편향·정보 접근 불평등·프라이버시 침해 등 복잡한 윤리적 도전을 야기하고 있음을 강조하며, 기존 디지털 헬스 리터러시 척도가 이러한 윤리적 차원을 포괄하지 못한다고 지적하였다. 예를 들어, DHLI는 신뢰성 평가와 개인정보 보호를 포함하더라도, AI 의료기술이 내리는 의사결정 과정에서 나타나는 공정성과 책임성, 투명성, 해명 가능성과 같은 핵심 윤리적 요소를 충분히 다루지 못한다. 이러한 한계는 AI 의료기술이 의료 형평성을 보장하지 못하거나 신뢰할 수 없는 의사결정을 초래할 가능성을 내포한다는 점에서 주목할 만하다. Lee et al.[22]은 COSMIN(Consensus-based Standards for the selection of health Measurement Instruments) 방법론을 활용하여 기존 디지털 헬스 리터러시 척도가 정보 소비 중심의 특성으로 인해 상호작용적이고 윤리적인 요구를 충분히 반영하지 못함을 확인하였다.

디지털 기술의 발전은 단순한 정보 활용 능력의 차원을 넘어 보다 심도 있는 윤리적 판단을 요구하게 되었으며, 이는 자연스럽게 디지털 윤리와 AI 윤리의 개념으로 연결된다. 디지털 윤리는 기술 사용의 책임성과 인간 중심적 접근을 강조하는 반면[23],[24], AI 윤리는 이러한 논의를 구체적으로 구현하여 공정성, 투명성, 책임성 등의 원칙을 중심으로 기술의 윤리적 활용 기준을 제시한다[25]. 특히, 의료 분야에서 AI 기술은 단순한 기술적 정확성을 넘어 높은 수준의 윤리적 고려를 요구하므로, 기존 디지털 헬스 리터러시 척도가 갖는 한계를 보완하기 위한 새로운 접근의 도입이 시급하다. 나아가, AI 의료기술 환경의 특수성을 반영하여 단순한 정보 접근 및 활용 능력을 넘어 공정성, 책임성, 투명성, 해명 가능성, 개인정보 보호 등 다양한 윤리적 요소를 포괄하는 개념의 확장이 필요하다. 이러한 확장은 알고리즘 편향과 차별적 결과를 방지하고, 사용자가 AI 의사결정 과정을 명확히 이해하며 그 결과에 대한 책임 소재를 분명히 하는 데 기여할 것으로 보인다.

이와 같이, AI 기반 헬스케어 기술이 제기하는 윤리적 쟁점과 기존 디지털 헬스 윤리 이론 및 척도의 한계를 종합적으로 고찰할 때, 단순 정보 활용 능력을 넘어 AI 의료기술의 공정성·투명성·책임성·개인정보 보호 등 윤리적 요소를 체계적으로 평가하고 교육할 수 있는 통합적 접근의 필요성이 대두된다. 이러한 통합적 프레임워크는 헬스 리터러시, 디지털 헬스 리터러시, AI 리터러시, 윤리 리터러시의 개념을 모두 포괄하여, AI 헬스케어 기술의 신뢰성과 사회적 수용성을 제고할 수 있는 기반을 마련할 것이다.

2-3 AI 헬스 윤리 리터러시: 개념적 정의와 구성요소

AI 기반 헬스케어 기술이 제공하는 혁신적 가능성을 온전히 이해하기 위해서는, 기존 디지털 헬스 리터러시의 범위를 넘어 AI 의료기술이 요구하는 윤리적 차원을 체계적으로 반영할 필요가 있다. 이에 따라 AI 헬스 윤리 리터러시는 단순한 정보 탐색과 활용 능력을 넘어, 공정성·책임성·투명성·개인정보 보호 등 다양한 윤리적 요소를 통합하여 사용자가 AI 기술의 신뢰성과 사회적 수용성을 높일 수 있는 역량을 배양하도록 지원하는 개념으로 정의된다. 즉, 이 개념은 AI 헬스케어 기술이 야기하는 복잡한 윤리적 문제들을 비판적으로 인식·평가하고, 그에 따른 적절한 윤리적 의사결정을 내릴 수 있는 실천적 역량을 의미한다.

이 개념은 헬스 리터러시, 디지털 헬스 리터러시, AI 리터러시, 그리고 윤리 리터러시라는 네 가지 기존 개념을 통합한 확장적 접근에 기초한다. 헬스 리터러시는 건강 정보를 찾고, 이해하며, 평가하여 적절한 의사결정을 내리는 능력을 의미한다[26]. 디지털 헬스 리터러시는 디지털 환경에서 정보를 탐색, 분석하고 문제를 해결하는 역량을 포함하며[13],[21], AI 리터러시는 AI 기술의 작동 원리와 한계를 이해하고 비판적으로 평가하는 능력을, 윤리 리터러시는 윤리적 문제를 식별·분석하여 공정하고 책임감 있는 결정을 내릴 수 있는 역량을 내포한다[27].

AI 헬스 윤리 리터러시의 핵심 구성요소는 기존 디지털 헬스 리터러시 척도와 구별되어, AI 헬스케어 기술에서 발생하는 윤리적 문제들을 효과적으로 다루기 위한 고유한 평가 요소들을 포함한다. 우선, “투명성”과 “해명 가능성”은 사용자가 AI 시스템의 의사결정 과정과 그 근거를 명확히 이해하고 비판적으로 평가할 수 있는 능력을 의미한다. 예를 들어, 사용자가 특정 진단이나 치료 권고의 데이터 근거를 파악하는 과정이 이에 해당한다. “공정성”과 “형평성”은 AI 시스템에 내재된 편향성을 인식하고, 이를 개선하기 위한 대안을 제시하거나 문제 개선을 요구할 수 있는 역량을 나타낸다. “책임성”은 AI 기술의 오류나 오작동 시 책임 소재를 명확히 파악하고 적절한 대응 조치를 취할 수 있는 능력을 강조하는데, 이는 의료 사고 발생 시 신속하고 체계적인 후속 조치로 이어진다.

“개인정보 보호”는 사용자가 자신의 민감한 건강 데이터가 AI 시스템에 의해 안전하게 관리되고 프라이버시가 침해되지 않도록 평가할 수 있는 능력을 포함하며, 이는 환자 신뢰의 기초로 작용한다. “위험 예방과 모니터링 능력”은 AI 헬스케어 기술과 관련된 잠재적 위험을 사전에 인식하고 예측하며, 필요 시 적절한 대처 전략을 제안할 수 있는 역량을 의미한다. 더불어, “다양성 존중”은 AI 기술이 특정 집단에 미칠 수 있는 편향적 영향을 평가하고 개선 대안을 모색하는 능력으로 설명할 수 있다. “자율성”은 사용자가 AI 의료기술을 활용하여 자신의 건강 관리와 관련된 의사결정을 독립적으로 내릴 수 있는 능력을 의미하며, 예를 들어 사용자가 AI의 추천 결

과와 자신의 의료적 판단을 비교·분석하는 과정이 자율성의 한 측면으로 볼 수 있다. 마지막으로, “신뢰”는 사용자가 AI 의료기술의 전반적 신뢰성을 기술적·윤리적 관점에서 평가하고 합리적 판단을 내릴 수 있는 능력을, “정보성”은 AI 시스템이 제공하는 정보의 정확성과 신뢰성을 평가하여 건강 관리 의사결정에 효과적으로 활용할 수 있는 역량을 나타낸다.

이와 같이, AI 헬스 윤리 리터러시는 다양한 윤리적 요소들을 체계적으로 통합함으로써, 사용자가 AI 헬스케어 기술이 초래하는 복잡한 윤리적 문제를 인식하고 비판적으로 평가할 수 있는 실천적 역량을 제공하는 데 중점을 둔다. 이러한 개념적 틀은 향후 AI 헬스케어 기술의 신뢰성과 사회적 수용성을 제고하기 위한 교육 및 평가 도구 개발의 이론적 기반으로 활용될 수 있다.

III. 연구방법

3-1 연구 설계 및 문헌검토를 통한 예비 문항 구성

본 연구는 AI 헬스케어 기술 사용 과정에서 발생할 수 있는 윤리적 문제를 인지하고 비판적으로 평가하며 책임감 있게 행동하는 능력을 “AI 헬스 윤리 리터러시”로 정의하고, 이를 측정할 수 있는 척도 개발을 목표로 하였다. 이를 위해 본 연구는 문헌검토를 통한 예비 문항 개발과 델파이 조사를 통한 문항 정교화의 두 단계로 구성된 연구 설계를 수립하였다.

우선, 헬스 리터러시, 디지털 헬스 리터러시, 윤리 리터러시, AI 리터러시 등 기존 척도의 구성 요인과 측정 항목을 분석하고, AI 윤리와 의료 윤리에 관한 주요 이론적 논의에서 제시된 핵심 원칙과 쟁점을 종합적으로 고찰하였다. 이러한 문헌 검토를 바탕으로, 본 연구는 AI 헬스 윤리 리터러시를 단순히 건강 정보를 탐색하고 활용하는 능력을 넘어, “AI 헬스케어 기술 사용 과정에서 발생할 수 있는 윤리적 문제를 비판적으로 평가하고 해결할 수 있는 사용자 중심의 역량”으로 정의하였다.

문헌 검토를 통해 도출된 AI 헬스 윤리 리터러시의 11가지 구성 요인과 예비 문항은 표 1에 제시하였다. 이러한 예비 문항은 AI 헬스 윤리 리터러시의 특성을 포괄적으로 반영하고자 설계되었으며, 이후 델파이 조사를 통해 정교화하였다.

3-2 델파이 조사 설계

델파이 조사는 특정 주제에 대한 전문가들의 합의를 도출하기 위해 고안된 체계적이고 반복적인 연구 방법으로, 전문가들의 익명성을 보장하고 특정 개인이나 집단의 의견 편향을 최소화하며, 객관적인 결과를 도출할 수 있다는 장점을 가진다[28],[29]. AI 헬스 윤리 리터러시는 AI·의료·윤리·사용자 경험 등 다양한 학문 분야가 결합된 복잡한 개념이기 때문

에, 다양한 분야의 전문가 참여를 통해 척도의 적합성과 완성도를 높이는 것이 필수적이다.

본 연구는 세 차례의 델파이 조사를 통해 AI 헬스 윤리 리터러시 척도의 구성 요소와 문항을 정교화하였다. 본 연구에 참여한 전문가는 다음과 같은 기준에 따라 선정하였다. 첫째, AI 윤리, 의료 윤리, 의료 정보학, 헬스 커뮤니케이션, 의료 법학, 뉴미디어 법 등 관련 분야에서 최소 5년 이상의 실무 또는 학술 연구 경험을 가진 전문가로 제한하였다. 둘째, 전문가의 연구 실적, 학술적 활동, 실무 경험의 다양성을 고려하여 균형 있게 선정하였다.

1차 조사는 AI 헬스 윤리 리터러시의 개념적 틀을 정립하고, 윤리 개념과 리터러시 역량 간의 연결 방식을 논의하는데 초점을 맞추었다. 2차 조사는 구성 요인 및 문항의 적절성, 명확성, 중요도를 평가하고, 수정·보완 작업을 수행하였다. 3차 조사는 최종 문항과 척도의 구조적 적합성을 검토하고 전문가들의 합의를 도출하는 데 목적을 두었다.

자료 수집은 2023년 11월부터 2024년 6월까지 약 8개월에 걸쳐 진행되었다. 각 조사 단계는 약 2~3주간 소요되었으며, 이메일을 통해 전문가들에게 참여를 요청한 후 서면 인터뷰 방식으로 의견을 수집하였다. 추가적인 답변이나 문항 수정 의견은 이메일로 회신받는 방식으로 자료를 수집하였다. 수집된 자료는 각 조사 단계에서 제시된 문항에 대한 전문가 의견과 제안을 종합적으로 검토하고, 5점 리커트 척도를 활용하여 문항의 적합성, 명확성, 중요성 등을 평가하여 최종적으로 합의된 문항과 구성 요소를 선정하였다.

총 15명의 전문가가 조사에 섭외되었으며, 단계별 참여 인원은 1차 조사 5명, 2차 조사 15명, 3차 조사 12명이었다. 1차 조사는 2023년 11월에 5명의 전문가를 대상으로 진행되었으며, AI 윤리 개념을 사용자 주관적 역량으로 구체화하고 윤리 인식을 평가 가능하도록 정의하는 데 초점을 맞추었다. 전문가 의견을 통해 지식 차원이 윤리적 인식 평가에 적합하지 않다는 결론이 도출되어, 지식 차원을 척도에서 제외하였다.

2차 조사는 2024년 4월에 15명의 전문가가 참여하여 진행되었다. 1차 조사에서 도출된 결과를 기반으로 재구성된 문항과 구성 요인을 평가하였다. 전문가들은 문항의 적절성, 명확성, 중요성을 검토하였으며, 각 구성 요인이 AI 헬스 윤리 리터러시의 핵심 요소로 적합한지를 중점적으로 논의하였다. 의견 수렴 과정은 이메일을 통해 진행되었으며, 이를 통해 문항의 수정 및 보완 작업이 이루어졌다.

3차 조사는 2024년 6월에 12명의 전문가가 참여하여 최종 문항과 구성 요인의 적합성을 평가하였다. 전문가들은 척도의 구조, 문항 수, 응답 형식 등 세부 사항을 검토하고 최종 합의를 도출하였다. 최종 단계에서도 서면 인터뷰와 이메일 교환 방식을 활용하여 전문가들의 의견을 심층적으로 반영하였다.

IV. 연구결과

4-1 1차 델파이 조사 결과

1차 델파이 조사는 AI 헬스 윤리 리터러시 척도 개발을 위한 기초 단계로, 5명의 전문가를 대상으로 수행되었다. 본 조사에서는 AI 헬스 윤리 리터러시의 개념을 정립하고, 윤리 개념과 이를 리터러시라는 사용자의 주관적 역량과 연결하는 방식을 중점적으로 논의하였다.

1) 윤리 개념에 대한 논의 및 측정 가능성 탐색

윤리의 본질과 측정 가능성에 대한 논의에서는 윤리가 도덕적 행동, 도덕적 판단, 도덕적 동기, 도덕적 인식 등 다차원적인 요소로 구성된다는 점을 재확인하였다[30],[31]. 전문가들은 이 중 “윤리적 인식”을 측정 가능하며, AI 헬스 윤리 리터러시의 핵심 구성 요소로 정의할 것을 제안하였다. “윤리적 판단은 민감성과 추론 능력을 통해 실천적으로 드러나야 한다(전문가A)”는 의견에 따라, AI 헬스케어 기술 사용과 관련된 다양한 윤리적 쟁점에 대한 사용자의 민감성을 “윤리 인식”으로 개념화하고, 이를 측정 가능한 형태로 구체화하고자 하였다.

또한, 전문가들은 윤리적 인식을 측정 가능한 형태로 구체화하는 과정에서, 이를 단순히 지식 차원으로 환원해서는 안 된다는 점을 강조하였다. “윤리는 주관적 사고와 실천적 역량을 포함하는 개념으로 접근해야 한다(전문가B)”라는 지적처럼, 윤리적 인식은 단순한 지식 습득을 넘어, 책임감 있게 행동할 수 있는 역량을 포괄하는 개념으로 이해되어야 한다는 것이다.

2) 인공지능 윤리 기준과 사용자 중심 평가

전문가들은 인간성을 중심으로, 인간 존엄성, 기술의 합목적성, 사회의 공공선 원칙에 따라 AI 헬스케어 기술을 평가해야 한다는 데 의견을 같이했다. 이를 위해 투명성, 안전성, 책임성, 인권 보장, 프라이버시 보호, 다양성 존중, 차별 금지, 공공성, 연대성, 데이터 관리, 공정성 등 인공지능 윤리 기준을 척도 개발에 반영해야 함을 강조하였다[32],[33]. 특히, 이러한 기준들을 사용자 관점에서 평가하기 위해, 사용자가 AI 헬스케어 기술 사용 과정에서 경험할 수 있는 구체적인 윤리적 상황과 이에 대한 인식과 태도를 측정하는 방안이 논의되었다. 한 전문가는 “윤리적 판단은 도덕적 민감성과 행동을 이끌어낼 수 있는 의사결정 역량으로 평가되어야 한다(전문가C)”라고 언급하며, 윤리적 인식이 구체적인 행동 및 의사결정으로 이어질 수 있는 실천적 역량 평가의 중요성을 강조하였다.

3) 지식 차원의 포함 여부에 대한 논의

당초 연구진은 AI 헬스 윤리 리터러시 척도를 인식, 태도, 지식의 세 가지 차원으로 구성하고자 하였다. 여기서 지식은 “디지털 인공지능 의료 서비스를 주체적이고 합리적으로 이용하기 위해 사용자 입장에서 알 필요가 있는 사실”로 정의되

었다. 이는 디지털 콘텐츠 소비 능숙도와 건강 정보를 구별할 수 있는 사고 능력 등을 포함하는 개념으로 제안되었다.

그러나 전문가들은 지식을 윤리적 인식 척도에 포함하는 것의 타당성과 현실적 가능성에 대해 다양한 의견을 제시하였다. 일부 전문가들은 지식의 포함이 사용자들이 AI 의료 서비스의 윤리적 측면을 얼마나 정확히 이해하고 있는지를 평가할 수 있다는 점에서 유용할 것이라고 평가하였다. 한 전문가가는 “지식을 측정함으로써 사용자의 인식이 실제 지식에 기반한 것인지, 아니면 잘못된 정보에 기반한 것인지를 구분할 수 있다”고 언급하였다. 그러나 다른 전문가들은 지식의 정의가 모호하며, AI 의료 서비스와 관련된 지식 수준은 개인의 배경에 따라 편차가 크기 때문에 평가의 객관성을 보장하기 어렵다는 점을 지적하였다. 한 전문가가는 “지식은 사실에 concept 대한 객관적이고 과학적인 판단을 할 수 있는 기본 요소이지만, 개인마다 편차가 크게 나타날 수 있다”며, 지식

을 척도에 포함하는 것의 한계를 지적하였다.

결론적으로, 전문가 다수는 윤리가 “단순히 알거나 배우는 개념이 아니라, 문제를 해결하고 책임 있는 선택을 내릴 수 있는 역량”이라는 점을 강조하며, 지식 차원을 척도에서 제외할 것을 제안하였다. 이에 본 연구는 기존의 디지털 리터러시[34], 헬스 리터러시[35], 디지털 헬스 리터러시[13],[21] 척도와 유사하게, AI 헬스 윤리 리터러시 척도를 주관적 신념과 실천적 역량 중심으로 구성하기로 결정하였다.

4) AI 헬스 윤리 리터러시 척도의 개념 정립

본 조사에서는 척도의 명칭과 관련하여 “AI 헬스 윤리 인식”과 “AI 헬스 윤리 리터러시” 두 가지 개념에 대한 심층적인 논의가 이루어졌다. 전문가들은 “인식”이 개인의 주관적인 신념 체계와 가치관을 포괄하는 더 넓은 개념이며, “리터러시”는 특정 영역에 대한 지식·기술·태도를 종합적으로 평가

표 1. AI 헬스 윤리 리터러시 척도의 구성 요소

Table 1. Components of AI health ethics literacy scale

Components	Initial concept	Revised concept
Transparency	The extent to which users understand the decision-making criteria and processes of AI medical services, and perceive the disclosure of such information as appropriate.	Clearly disclose and explain AI system data processing, algorithms, decision-making processes, and outcomes.
Communication	The extent to which users evaluate the clarity and efficiency of communication with AI medical services, and perceive their understanding of the provided information.	Provide clear and accurate information considering diverse user characteristics, ensuring respectful and effective communication.
Fairness	The extent to which users perceive that AI medical services provide fair treatment to all patients without bias.	Prevent discrimination based on individual characteristics and ensure fair outcomes grounded in scientific data.
Risk Prevention & Monitoring	The extent to which users recognize the AI medical service's ability to prevent and monitor potential risks, and how important they consider such risk management procedures.	Predict, notify, prevent, and manage risks associated with AI services while providing feedback.
Respect for Diversity	The extent to which users perceive that AI medical services respect patient diversity and offer personalized treatments.	Ensure respect for individual and group characteristics throughout service provision, guaranteeing non-discriminatory accessibility.
Safety	The extent to which users perceive the safety of AI medical services and feel secure using such systems.	Prevent errors and security breaches while ensuring a safe service environment with a responsive system for issue resolution.
Accountability	The extent to which users recognize the accountability of AI medical service providers and their expectations regarding response measures in case of issues.	Clearly define responsibilities in service design and operation, while ensuring proper responses and remedies in case of problems.
Privacy Protection	The extent to which users perceive their personal data as well-protected within AI medical services and their level of trust in data security.	Guarantee security and legal compliance throughout data collection, processing, storage, and disposal, while ensuring user control over personal data.
Autonomy	The extent to which users feel free to express their opinions and recognize their decision-making rights in AI medical service usage.	Ensure that users can independently decide whether to use the service and accept its results.
Trust	The extent to which users trust AI medical services and their providers, and how this trust impacts their experience.	Build mutual trust between users and providers based on the reliability of AI services.
Information Quality	The extent to which users perceive the quality and timeliness of the information provided by AI medical services, and its usefulness in health management and decision-making.	Provide timely, high-quality information while maintaining accuracy, relevance, and reliability.

하는 실천적 역량에 가깝다는 점에 주목하였다. 논의 결과, 본 연구는 AI 헬스케어 기술 사용과 관련된 윤리적 쟁점을 비판적으로 성찰하고, 책임감 있게 행동하는 실천적 역량 함양을 목표로 한다는 점에서 “AI 헬스 윤리 리터러시”를 척도의 명칭으로 최종 확정하였다.

4-2 2차 델파이 조사 결과

2차 델파이 조사는 1차 조사 결과를 바탕으로 재구성된 AI 헬스 윤리 리터러시 척도의 구성 요소에 대해 전문가들의 평가와 추가 의견을 수렴하기 위해 15명의 전문가를 대상으로 실시되었다. 이번 조사의 주요 목적은 각 구성 요소의 중요도를 평가하고, 개념적 정의를 검토하기 위한 질적 의견을 수집하는 데 있었다. 전문가들은 각 항목의 중요도를 5점 척도로 평가하였으며, 개념적 정의와 추가적인 수정 방향에 대한 의견을 서면으로 제시하였다.

1) 중요도 평가 결과

전문가들은 AI 헬스 윤리 리터러시의 11개 구성 요소(투명성, 의사소통, 공정성, 위험 예방 및 모니터링, 다양성 존중, 안전성, 책임성, 개인정보 보호, 자율성, 신뢰, 정보성)에 대해 평균적으로 높은 중요도를 부여하였다(그림 1).



그림 1. AI 헬스 윤리 리터러시 구성 요소의 중요도 평가 결과
Fig. 1. Importance evaluation results of AI health ethics literacy components

분석 결과, 안전성(4.75 ± 0.71 점)과 책임성(4.75 ± 0.46 점)은 가장 높은 중요도를 기록하여 전문가들이 해당 항목을 AI 헬스 윤리 리터러시에서 가장 핵심적인 요소로 인식하고 있음을 보여주었다. 개인정보 보호(4.63 ± 0.74 점), 공정성(4.50 ± 0.76 점), 위험 예방 및 모니터링(4.38 ± 0.74 점) 항목 역시 높은 중요도를 나타냈다. 반면, 정보성(3.50 ± 1.20 점)은 가장 낮은 평균 점수를 기록했으며, 전문가 간 의견 차이가

비교적 컸음을 확인할 수 있었다.

2) 구성 항목별 전문가 의견과 수정 방향

2차 델파이 조사를 통해, AI 헬스 윤리 리터러시 척도의 구성 요소를 구체화하고, 전문가들의 질적 의견을 반영하여 각 항목에 대한 정의를 체계적으로 수정하였다(표 1 참조). 특히, 전문가들의 의견은 기존 구성 요소 간 중복성을 줄이고, 각 항목의 타당성과 실효성을 높이는 방향으로 정의를 구체화하는 데 기여하였다. 이러한 결과는 3차 델파이 조사에서 척도를 정교화하고 최종 문항을 확정하는 데 중요한 근거로 활용되었다.

4-3 3차 델파이 조사 결과

3차 델파이 조사는 AI 헬스 윤리 리터러시 척도의 문항 적절성을 평가하고, 문항의 신뢰도를 분석하여 최종 척도를 정교화하기 위해 진행되었다. 12명의 전문가가 참여하였으며, 11개의 구성 요소와 이에 해당하는 총 33개 문항을 대상으로 평가가 이루어졌다. 전문가들은 각 문항의 적절성을 5점 척도로 평가하였고, 추가적으로 문항 수정 방향에 대한 의견을 주관식으로 제시하였다.

1) 구성 요소별 신뢰도 분석: 척도의 일관성 검증

크로바흐 알파(Cronbach's Alpha) 계수를 통해 각 구성 요소의 내부 신뢰도를 검토한 결과, 신뢰도 계수는 .670에서 .857 사이로 나타났다. 의사소통 항목은 알파 값 .857로 가장 높은 신뢰도를 기록하며, 문항들이 구성 요소를 일관되게 설명하고 있음을 확인하였다. 반면, 안전성 항목은 .670으로 비교적 낮은 신뢰도를 보여, 문항 재구성 또는 세부 문항의 보완이 필요하다는 시사점을 제공하였다. 책임성(.777), 개인정보 보호(.755), 공정성(.740) 등은 신뢰도가 안정적으로 나타나, 해당 구성 요소가 문항 간 높은 일관성을 유지하고 있음을 보여준다(표 2).

2) 문항별 적절성 평가: 측정 문항의 타당성 검토

문항 적절성 평가의 평균 점수는 2.88에서 4.88 사이로 나타났다. 이를 통해 각 구성 요소와 문항의 상관성을 보다 구체적으로 검토할 수 있었다. 가장 높은 점수를 기록한 문항은 위험 예방 및 모니터링 항목의 “AI 의료 시스템 사용 중 발생할 수 있는 위험에 대해 사전에 안내를 받았다고 생각하십니까?”로, 4.88 ± 0.35 를 기록하며 전문가들로부터 높은 적절성을 보였다. 반면, 신뢰 항목의 “디지털 의료 서비스의 신뢰성을 높이기 위해 가장 중요하다고 생각하는 요소는 무엇입니까?” 문항은 2.88 ± 1.13 로 가장 낮은 점수를 보였으며, 구체성이 부족하다는 의견과 함께 수정의 필요성이 제기되었다.

전문가들은 각 문항의 적절성과 관련하여 문항의 명확성, 이해 용이성, 측정 대상과의 연관성을 강화할 것을 제안하였다. 투명성 항목에서는 데이터 처리와 의사결정 과정의 투명

표 2. AI 헬스 윤리 리터러시 척도 구성 문항

Table 2. AI health ethics literacy scale items

Components	Items	Mean±Standard Deviation	Cronbach's alpha
Transparency	To what extent do you think AI medical services clearly share their decision-making processes?	3.75±0.89	.709
	Do you feel that you receive sufficient information to understand how AI medical systems operate?	3.50±1.20	
	Are you satisfied with the explanations provided on how AI medical systems generate diagnoses or treatment recommendations?	3.75±1.58	
Communication	Do you find communication with AI medical service providers clear and easy to understand?	4.75±0.46	.857
	Do you feel that you can request additional explanations regarding AI medical service recommendations?	3.13±1.13	
	When you have questions about AI medical service decisions, do you feel that there are sufficient channels to obtain further information?	4.00±0.93	
Fairness	Do you think AI medical services treat all patients fairly?	4.75±0.46	.750
	Do you perceive AI-based diagnoses as unbiased and fair across different patient groups?	4.38±0.92	
	Do you believe that the data used in AI medical services adequately represent diverse patient populations?	3.13±0.64	
Risk Prevention & Monitoring	Do you feel that AI medical services take sufficient measures to identify and respond to potential risks?	4.00±1.20	.825
	Have you received prior guidance on potential risks associated with using AI medical systems?	4.88±0.35	
	Do you consider the risk management and monitoring procedures in AI medical services to be effective?	3.75±0.71	
Respect for Diversity	Do you feel that AI medical services take cultural and social backgrounds into account?	4.25±0.71	.800
	Do you think AI-based medical services recognize and appropriately respond to individual differences?	4.50±0.76	
	Do you believe that AI medical systems sufficiently consider the needs and circumstances of diverse patients?	4.13±0.64	
Safety	Do you feel that using AI medical services ensures patient safety?	3.75±1.17	.670
	In case of safety-related issues in AI medical services, do you believe there are adequate procedures in place to resolve them?	4.13±1.13	
	How would you assess the safety of AI-based medical services?	3.13±1.13	
Accountability	Do you think it is clear who is responsible for issues arising in AI medical services?	4.38±0.74	.740
	Do you believe that there are sufficient procedures to address errors in AI medical systems?	4.63±0.52	
	Do you think clarifying the responsibilities of AI medical service providers improves service quality?	3.25±0.89	
Privacy Protection	Are you satisfied with the level of personal data protection provided by AI medical services?	3.88±1.13	.777
	Do you find the privacy policy clearly disclosed and easy to understand?	4.38±0.74	
	Do you have concerns about the misuse of your personal data?	3.63±1.19	
Autonomy	Do you feel that you have the freedom to make your own health-related decisions while using AI medical services?	4.13±0.99	.755
	If you disagree with AI medical service recommendations, do you find it easy to express your opinion?	3.88±0.84	
	Do you think personal choice should take precedence over AI recommendations in medical decision-making?	3.63±1.19	
Trust	Do you trust AI medical services and feel comfortable following their recommendations?	4.25±0.89	.685
	How do you perceive the expertise and reliability of AI medical service providers?	3.38±0.74	
	What do you consider the most important factor in increasing trust in digital medical services?	2.88±1.13	
Information Quality	Do you find the information provided by AI medical services useful and reliable?	4.00±1.07	.800
	Are you satisfied with the accuracy and up-to-date nature of the medical information provided?	3.75±0.17	
	Do you believe that the information received from AI medical services helps you manage your health and make informed decisions?	4.00±1.41	

표 3. 최종 AI 헬스 윤리 리터러시 척도 문항

Table 3. Final AI health ethics literacy scale items

Components	Definitions	Items
Transparency in Decision-Making	The ability to recognize the transparency of AI medical technology in data processing and decision-making, and to assess the accessibility and explainability of information.	Do you feel that the decision-making process of AI healthcare services is clearly explained?
		Do you think that the data and criteria used in AI healthcare service decisions are transparently disclosed?
		Do you feel that information related to AI healthcare services is easily accessible?
Communication with Healthcare Providers	The ability to communicate clearly and effectively with AI medical technology providers, understand, and utilize necessary information.	Do you find the information provided by AI healthcare service providers to be clear and easy to understand?
		Do you feel that there are communication channels available to ask questions or provide feedback about AI healthcare services?
		Do you think communication with healthcare providers is effective and smooth?
Fairness in Decision-Making	The ability to evaluate whether AI medical technology is applied fairly to all users and operates objectively without bias.	Do you believe AI healthcare services operate fairly without bias against specific groups?
		Do you feel that AI healthcare service outcomes are applied equally to all users?
		Do you think AI healthcare services reflect the needs and demands of diverse users?
Risk Management and Response	The ability to recognize potential risks associated with AI medical technology usage and understand the systems for prevention and response.	Do you believe that AI healthcare services proactively identify potential risks and provide preventive measures?
		Have you received clear response procedures for issues that arise in AI healthcare services?
		Do you think AI healthcare services appropriately handle emergency situations?
Respect for Patient Diversity	The ability to assess whether AI medical technology respects the cultural, social, and personal backgrounds of diverse patients and provides inclusive services.	Do you feel that AI healthcare services respect users' cultural and social backgrounds?
		Do you believe AI healthcare services adequately reflect individual needs?
		Do you think AI healthcare services sufficiently consider the diverse circumstances of users?
System Safety	The ability to evaluate the stability and security of AI medical technology and understand how to effectively address issues when they arise.	Do you feel that AI healthcare services ensure user safety?
		Do you trust the security and safety of AI healthcare services?
		Do you believe there are sufficient procedures in place to resolve system errors when they occur?
Accountability and Responsibility	The ability to clearly understand accountability in case of issues related to AI medical technology and assess the problem-resolution process.	Do you clearly understand who is responsible when issues arise in AI healthcare services?
		Do you think clear procedures are in place to address errors in AI healthcare services?
		Do you feel that accountability is well-defined in AI healthcare services?
Autonomy	The ability to understand and exercise the right to make independent health management decisions while using AI medical technology.	Do you feel that you can make your own health-related decisions while using AI healthcare services?
		If you disagree with AI healthcare service recommendations, do you feel free to reject them?
		Do you believe AI healthcare services fully guarantee user autonomy?
Trust in AI Healthcare Services	The ability to trust the information and recommendations provided by AI medical technology and make informed decisions based on them.	Do you trust the recommendations provided by AI healthcare services?
		Do you think AI healthcare service providers offer reliable information?
		Do you feel that AI healthcare services are generally trustworthy?
Quality of Information Provided	The ability to evaluate the usefulness, diversity, and accuracy of the information provided by AI medical technology.	Do you think AI healthcare services provide sufficient information to meet your needs?
		Do you believe AI healthcare services offer a diverse range of healthcare-related information?
		Do you find the information provided by AI healthcare services to be sufficiently useful?
Self-Efficacy	The belief in one's ability to effectively utilize AI medical technology for personal health management.	Do you believe that AI healthcare services help you manage your health more effectively?
		Do you feel that using AI healthcare services improves your ability to make medical decisions?
		Do you think the information provided by AI healthcare services helps you effectively address health concerns?
Privacy Protection	The ability to assess whether AI medical technology adequately protects personal data and ensures user control over personal information management.	Do you believe AI healthcare services protect personal data securely?
		Do you feel that privacy policies are clearly disclosed?
		Do you have no concerns about the improper use of your personal data?

성을 강조하면서 접근성 관련 문항을 추가할 필요가 제기되었다. 공정성과 다양성 존중은 정의와 문항 간 중복을 최소화하고, 소비자가 실제로 판단할 수 있는 질문으로 수정할 것을 권고하였다. 위험 예방 및 모니터링 문항은 소비자가 위험 관리 절차를 직접 평가하기 어렵다는 점을 반영해 대처 방안과 사전 안내에 중점을 둔 질문으로 변경해야 한다는 의견이 있었다. 의사소통과 자율성은 문항 간 내용의 경계가 불분명하다는 지적을 바탕으로, 권고 사항에 대한 추가 설명과 의견 제시 문항을 자율성으로 이동하도록 조정해야 했다. 안전성은 단순한 안정성을 넘어 보안성과 회복 탄력성을 강조하는 방향으로 문항을 수정해야 한다는 제안이 있었으며, 책임성 항목은 사전적 책임과 사후적 책임을 명확히 구분하여 문항을 재구성할 필요가 있다고 평가되었다. 신뢰와 정보성 항목에서는 중복된 질문을 제거하고, 신뢰성과 최신성이 사용자 경험에 미치는 영향을 명확히 평가할 수 있는 문항으로 재구성해야 한다는 의견이 제시되었다. 이러한 의견은 최종 척도의 수정과 재구성에 중요한 근거로 작용하였다.

3) 수정된 AI 헬스 윤리 리터러시 척도

수정된 AI 헬스 윤리 리터러시 척도는 기존 11개 항목에서 12개 항목으로 확장되었으며, 전문가들의 질적 의견을 반영하여 각 항목의 조작성 정의와 측정 문항이 체계적으로 재구성되었다. 주요 변경 사항은 전문가들이 제안한 항목 간 중복 해소, 정의의 명확화, 측정 문항의 구체성을 강화하는 데 중점을 두었다. 특히, 새로운 요인인 효능감(Self-Efficacy)이 추가되었으며, 이는 사용자가 AI 의료 서비스를 사용할 때 자신의 의료 결정 능력과 건강 관리 능력이 향상되었다고 느끼는지를 측정하도록 설계되었다. 이는 사용자의 심리적 경험과 결과적 효능을 평가하기 위한 항목으로, 기존 척도의 내용적 범위를 확장하였다.

전문가 의견을 반영한 주요 수정 내용은 다음과 같다. 투명성(Transparency) 항목은 기존의 “결정 과정의 대등한 참여”를 공정성(Fairness)으로 이동하고, 대신 “정보 접근 용이성”을 추가하여 정의를 보완하였다. 의사소통(Communication)은 주로 정보 전달의 이해 용이성과 상호작용의 명확성을 강조하며, 일부 문항은 자율성(Autonomy) 항목으로 이동되었다. 공정성(Fairness)은 기존 다양성 존중(Bias Exclusion)과 중복된 요소를 통합하고, 객관적 데이터에 기반한 공정성을 강조하였다. 또한, 위험 관리 및 대응(Risk Management and Response) 항목은 “위험 예방”과 “위험 대응”을 모두 포함하여 정의를 확대하고, 사용자 체감 경험을 중심으로 문항을 재구성하였다. 전문가들은 일부 문항이 복합적인 질문을 포함하거나 사용자에게 응답하기 어려운 내용을 포함하고 있다는 점을 지적하였으며, 이에 따라 문항의 표현을 단순화하고, 응답 가능성을 높이도록 수정하였다. 이와 같은 수정은 AI 의료 서비스의 윤리적 측면을 사용자 관점에서 더욱 정교하게 평가할 수 있도록 하며, 척도의 타당성과 신뢰성을 높이는 중요한 조치로 작용하였다. 최종 문항은 표 3에 제시하였다.

V. 결론 및 논의

본 연구는 AI 헬스 윤리 리터러시(AI Health Ethics Literacy)의 개념을 정립하고, 윤리적 인식을 평가할 수 있는 초기 항목을 도출하는 것을 목표로 하였다. 이를 위해 델파이 기법을 활용하여 AI, 의료, 윤리 분야의 전문가 집단의 의견을 수렴한 결과, 기존 디지털 헬스 리터러시 척도의 한계를 보완하고 AI 의료기술의 윤리적 수용성을 평가할 수 있는 12개의 핵심 요소를 도출하였다. 기존 디지털 헬스 리터러시 연구들이 주로 정보의 접근성, 탐색 및 활용 능력을 중심으로 개인의 역량을 평가하는 데 머물렀다면, 본 연구는 사용자가 AI 헬스케어 기술 사용 과정에서 직면할 수 있는 윤리적 문제를 비판적으로 인지하고 책임 있게 대처하는 역량(윤리적 민감성, 책임성 등)을 구체화하여, AI 헬스케어 기술 평가의 새로운 틀을 제시하였다.

이러한 연구 결과는 여러 측면에서 학문적 기여를 제공한다. 첫째, 기존의 디지털 리터러시 및 헬스 리터러시 척도가 주로 정보 탐색 및 활용 능력에 초점을 맞춘 반면, 본 연구는 AI 의료기술과 관련된 윤리적 쟁점을 통합적으로 평가할 수 있는 개념적 틀을 마련하였다. 둘째, 공정성, 투명성, 책임성과 같은 윤리적 원칙을 구조화하고, 이를 사용자 관점에서 평가할 수 있도록 주요 항목을 도출함으로써, AI 의료기술의 신뢰성과 윤리적 수용성을 검토할 수 있는 기초 자료를 제공하였다. 셋째, 효능감(self-efficacy)과 같은 기존 연구에서 상대적으로 간과되었던 요인을 포함하여, 사용자가 AI 의료기술을 활용할 때 느끼는 심리적 자율성과 윤리적 판단 역량을 평가할 수 있는 새로운 차원을 제시하였다.

이러한 학문적 기여와 더불어, 본 연구의 결과는 실무적 측면에서도 중요한 의미를 가진다. 첫째, AI 헬스 윤리 리터러시 평가 항목을 통해 한국 사용자들의 AI 의료 윤리 인식 수준과 취약성을 진단할 수 있다. 이를 통해 의료 기술 도입 과정에서 발생할 수 있는 윤리적 문제를 선제적으로 파악하고 대응하는 데 활용될 수 있다. 둘째, AI 의료 윤리 교육 및 가이드라인 개발의 기초 자료로 활용될 수 있으며, 사용자 중심의 윤리 교육과 정책 설계에 기여할 수 있다. 의료 서비스 제공자들에게는 알고리즘의 투명성, 데이터 처리 절차, 사용자 의견 반영 등의 윤리적 기준을 강화하여 서비스 품질을 개선하는 데 실질적인 가이드라인을 제공할 수 있다. 더 나아가, 본 연구는 한국 사회의 윤리적 요구를 반영한 AI 의료 윤리 담론을 형성하고, 건강한 AI 의료 생태계를 조성하는 데 기여할 것으로 기대된다.

한편, 본 연구는 델파이 기법을 활용하여 전문가 합의를 도출하는 방식으로 진행되었으나 일반 사용자 집단에 대한 정량적 검증이 이루어지지 않았다는 한계를 가진다. 따라서 향후 연구에서는 일반 사용자, 의료인, 정책 입안자 등 다양한 이해관계자를 대상으로 한 실증 연구를 수행하여 AI 헬스 윤리 리터러시 척도의 신뢰성과 타당성을 검증할 필요가 있다.

구체적으로, 설문조사 및 심층 인터뷰를 활용하여 척도의 구조적 타당성을 평가하고, 실제 디지털 헬스 기술 수용 과정에서 윤리적 인식이 어떻게 작용하는지를 실증적으로 분석하는 연구가 요구된다. 또한, 본 연구는 한국의 의료 환경과 사회문화적 맥락에 초점을 맞추었으므로, 타 국가 및 문화적 배경에서 AI 헬스 윤리 리터러시의 개념과 평가 항목이 어떻게 적용될 수 있는지 검토하는 국제적 연구가 요구된다. 마지막으로, AI 의료기술이 지속적으로 발전함에 따라 새로운 윤리적 쟁점이 등장할 가능성이 높으므로, 본 연구에서 도출된 평가 항목은 정기적으로 검토 및 보완될 필요가 있다.

한국 사회는 개인정보 보호와 의료 정보의 민감성에 대한 높은 인식으로 인해 AI 헬스케어 기술과 관련된 윤리적 문제에 민감하게 반응한다. 또한, 고령화와 만성질환 증가로 인해 AI 기반 헬스케어 기술이 의료 시스템의 부담을 완화할 중요한 대안으로 주목받고 있다[36]. 그러나 기술적 완성도만으로는 충분하지 않으며, 데이터의 투명성, 알고리즘의 공정성, 책임성 등 윤리적 기준이 명확히 마련되어야 AI 헬스케어 기술이 신뢰받을 수 있다. 이에 따라, 본 연구에서 도출한 AI 헬스 윤리 리터러시 평가 항목은 기술 설계 단계부터 윤리적 고려를 반영하고, 사용자들이 윤리적 문제를 인식하고 적절히 대응할 수 있도록 지원하는 중요한 역할을 할 것이다.

감사의 글

이 논문은 2022년 대한민국 교육부와 한국연구재단의 일 반공동연구지원사업의 지원을 받아 수행된 연구임(NRF-2022S1A5A2A03055583).

참고문헌

- [1] T. Davenport and R. Kalakota, "The Potential for Artificial Intelligence in Healthcare," *Future Healthcare Journal*, Vol. 6, No. 2, pp. 94-98, June 2019. <https://doi.org/10.7861/futurechosp.6-2-94>
- [2] H. Zhao, J. Cao, J. Xie, W. Liao, Y. Lei, H. Cao ... & C. Bowen, "Wearable Sensors and Features for Diagnosis of Neurodegenerative Diseases: A Systematic Review," *Digital Health*, 9, 20552076231173569, 2023. <https://doi.org/10.1177/205520762311735>
- [3] W. Salah Eldin and A. Kaboudan, "AI-Driven Medical Imaging Platform: Advancements in Image Analysis and Healthcare Diagnosis," *Journal of the ACS Advances in Computer Science*, Vol. 14, No. 1, 2023. <https://doi.org/10.21608/asc.2023.328064>
- [4] E. Vayena, A. Blasimme, and I. G. Cohen, "Machine Learning in Medicine: Addressing Ethical Challenges," *PLoS Medicine*, Vol. 15, No. 11, e1002689, November 2018. <https://doi.org/10.1371/journal.pmed.1002689>
- [5] S. Gerke, T. Minssen, and G. Cohen, Ethical and Legal Challenges of Artificial Intelligence-Driven Healthcare, in *Artificial Intelligence in Healthcare*, London, UK: Academic Press, ch. 12, pp. 295-336, 2020. <https://doi.org/10.1016/B978-0-12-818438-7.00012-5>
- [6] Z. Obermeyer, B. Powers, C. Vogeli, and S. Mullainathan, "Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations," *Science*, Vol. 366, No. 6464, pp. 447-453, October 2019. <https://doi.org/10.1126/science.aax2342>
- [7] I. G. Cohen, R. Amarasingham, A. Shah, B. Xie, and B. Lo, "The Legal and Ethical Concerns that Arise from Using Complex Predictive Analytics in Health Care," *Health Affairs*, Vol. 33, No. 7, pp. 1139-1147, July 2014. <https://doi.org/10.1377/hlthaff.2014.0048>
- [8] J. Morley, C. C. V. Machado, C. Burr, J. Cows, I. Joshi, M. Taddeo, and L. Floridi, "The Ethics of AI in Health Care: A Mapping Review," *Social Science & Medicine*, Vol. 260, 113172, September 2020. <https://doi.org/10.1016/j.socscimed.2020.113172>
- [9] N. R. Möllmann, M. Mirbabaie, and S. Stieglitz, "Is It Alright to Use Artificial Intelligence in Digital Health? A Systematic Literature Review on Ethical Considerations," *Health Informatics Journal*, Vol. 27, No. 4, p. 14604582211052391, 2021. <https://doi.org/10.1177/14604582211052391>
- [10] A. J. Larrazabal, N. Nieto, V. Peterson, D. H. Milone, and E. Ferrante, "Gender Imbalance in Medical Imaging Datasets Produces Biased Classifiers for Computer-Aided Diagnosis," *Proceedings of the National Academy of Sciences*, Vol. 117, No. 23, pp. 12592-12594, June 2020. <https://doi.org/10.1073/pnas.1919012117>
- [11] O. Norgaard, D. Furstrand, L. Klokke, A. Karnoe, R. Batterham, L. Kayser, and R. H. Osborne, "The e-Health Literacy Framework," *Knowledge Management & e-Learning*, Vol. 7, No. 4, pp. 522-540, 2015. <https://doi.org/10.34105/j.kmel.2015.07.035>
- [12] C. D. Norman and H. A. Skinner, "eHealth Literacy: Essential Skills for Consumer Health in a Networked World," *Journal of Medical Internet Research*, Vol. 8, No. 2, e9, 2006. <https://doi.org/10.2196/jmir.8.2.e9>
- [13] R. van der Vaart and C. Drossaert, "Development of the Digital Health Literacy Instrument: Measuring a Broad Spectrum of Health 1.0 and Health 2.0 Skills," *Journal of Medical Internet Research*, Vol. 19, No. 1, e27, January 2017. <https://doi.org/10.2196/jmir.6709>

- [14] S. Hongladarom and J. Bandasak, "Non-Western AI Ethics Guidelines: Implications for Intercultural Ethics of Technology," *AI & Society*, Vol. 39, No. 4, pp. 2019-2032, 2024. <https://doi.org/10.1007/s00146-023-01665-6>
- [15] J. H. Kim, H. S. Jung, M. H. Park, S. H. Lee, H. Lee, Y. Kim, and D. Nan, "Exploring Cultural Differences of Public Perception of Artificial Intelligence via Big Data Approach," in *Proceedings of the International Conference on Human-Computer Interaction*, Cham: Springer International Publishing, June 2022, pp. 427-432.
- [16] J. Amann, A. Blasimme, E. Vayena, D. Frey, V. I. Madai, and Precise4Q Consortium, "Explainability for Artificial Intelligence in Healthcare: A Multidisciplinary Perspective," *BMC Medical Informatics and Decision Making*, Vol. 20, No. 1, pp. 1-9, 2020. <https://doi.org/10.1186/s12911-020-01332-6>
- [17] K. H. Keskinbora, "Medical Ethics Considerations on Artificial Intelligence," *Journal of Clinical Neuroscience*, Vol. 64, pp. 277-282, June 2019. <https://doi.org/10.1016/j.jocn.2019.03.001>
- [18] A. Singhal, N. Neveditsin, H. Tanveer, and V. Mago, "Toward Fairness, Accountability, Transparency, and Ethics in AI for Social Media and Health Care: Scoping Review," *JMIR Medical Informatics*, Vol. 12, No. 1, e50048, 2024. <https://doi.org/10.2196/50048>
- [19] A. S. Albahri, A. M. Duham, M. A. Fadhel, A. Alnoor, N. S. Baqer, L. Alzubaidi, ... and M. Deveci, "A Systematic Review of Trustworthy and Explainable Artificial Intelligence in Healthcare: Assessment of Quality, Bias Risk, and Data Fusion," *Information Fusion*, Vol. 96, pp. 156-191, 2023. <https://doi.org/10.1016/j.inffus.2023.03.008>
- [20] K. Siau and W. Wang, "Artificial Intelligence (AI) Ethics: Ethics of AI and Ethical AI," *Journal of Database Management*, Vol. 31, No. 2, pp. 74-87, 2020. <https://doi.org/10.4018/jdm.2020040105>
- [21] C. D. Norman and H. A. Skinner, "eHealth Literacy: Essential Skills for Consumer Health in a Networked World," *Journal of Medical Internet Research*, Vol. 8, No. 2, e9, 2006. <https://doi.org/10.2196/Jmir.8.2.E9>
- [22] J. Lee, E. H. Lee, and D. Chae, "eHealth Literacy Instruments: Systematic Review of Measurement Properties," *Journal of Medical Internet Research*, Vol. 23, No. 11, e30644, 2021. <https://doi.org/10.2196/30644>
- [23] L. Floridi, *The Ethics of Information*, Oxford, UK: Oxford University Press, 2013.
- [24] P. Solanki, J. Grundy, and W. Hussain, "Operationalising Ethics in Artificial Intelligence for Healthcare: A Framework for AI Developers," *AI and Ethics*, Vol. 3, No. 1, pp. 223-240, 2023. <https://doi.org/10.1007/s43681-022-00195-z>
- [25] H. S. Lee and H. D. Cheon, "The Principle of Explicability in AI Ethics," *Study of Humanities*, Vol. 35, pp. 37-63, June 2021. <http://doi.org/10.31323/sh.2021.06.35.02>
- [26] D. Nutbeam, "Health Literacy as a Public Health Goal: A Challenge for Contemporary Health Education and Communication Strategies into the 21st Century," *Health Promotion International*, Vol. 15, No. 3, pp. 259-267, September 2000. <https://doi.org/10.1093/heapro/15.3.259>
- [27] K. Murphy, E. Di Ruggiero, R. Upshur, D. J. Willison, N. Malhotra, J. C. Cai, ... and J. Gibson, "Artificial Intelligence for Good Health: A Scoping Review of the Ethics Literature," *BMC Medical Ethics*, Vol. 22, No. 1, pp. 1-17, 2021. <https://doi.org/10.1186/s12910-021-00577-8>
- [28] C.-C. Hsu and B. A. Sandford, "The Delphi Technique: Making Sense of Consensus," *Practical Assessment, Research, and Evaluation*, Vol. 12, No. 1, 10, 2007. <https://doi.org/10.7275/pdz9-th90>
- [29] C. Okoli and S. D. Pawlowski, "The Delphi Method as a Research Tool: An Example, Design Considerations and Applications," *Information & Management*, Vol. 42, No. 1, pp. 15-29, December 2004. <https://doi.org/10.1016/j.im.2003.11.002>
- [30] J. R. Rest, *Moral Development: Advances in Research and Theory*, New York, NY: Praeger, 1983.
- [31] L. J. Walker and R. C. Pitts, "Naturalistic Conceptions of Moral Maturity," *Developmental Psychology*, Vol. 34, No. 3, pp. 403-419, 1998. <https://doi.org/10.1037/0012-1649.34.3.403>
- [32] European Commission, *Ethics Guidelines for Trustworthy AI*, Publications office of the European Union, Brussels, Belgium, April 2019.
- [33] IEEE (Institute of Electrical and Electronics Engineers) Standards Association, *Ethically Aligned Design: A Vision for Prioritizing Human Well-Being with Autonomous and Intelligent Systems*, Author, Piscataway: NJ, 2019.
- [34] P. Gilster, *Digital Literacy*, New York, NY: Wiley Computer Publishing, 1997.
- [35] K. Sørensen, S. Van Den Broucke, J. Fullam, G. Doyle, J. Pelikan, Z. Slonska, and H. Brand, "Health Literacy and Public Health: A Systematic Review and Integration of Definitions and Models," *BMC Public Health*, Vol. 12, 80, January 2012. <https://doi.org/10.1186/1471-2458-12-80>
- [36] H. Song, J. Kim, and T. Kim, "Artificial Intelligence and Health Communication," *Journal of Communication Research*, Vol. 57, No. 3, pp. 196-238, 2020. <https://doi.org/10.22174/jcr.2020.57.3.196>



이하나(Hannah Lee)

2021년 : 이화여자대학교 커뮤니케이션·미디어학부 (언론학 박사)

2021년~현 재: 이화여자대학교 연령통합고령사회연구소 연구교수

※ 관심분야 : 모바일 헬스(Mobile Health), HCI(Human Computer Interaction) 등



엄주희(Juhee Eom)

2013년 : 연세대학교 법학과 (법학 박사)

2020년~현 재: 건국대학교 융합인재학부 교수

※ 관심분야 : 헌법과 융합법, 뇌신경법학(Neurolaw), 인공지능과 인권(Artificial Intelligence and Human Rights) 등