

컴퓨터교육학회 논문지 2026년 제29권 제8호
<https://doi.org/10.32431/kace.2026.29.8.010>



AI 시대의 핵심 역량에 대한 탐구: 컴퓨팅사고력에서 AI 맥락 다관점 메타인지로의 이동

Exploring Core Competencies in the Age of AI: A Shift from Computational Thinking to AI-Situated Perspectival Metacognition

유인환[†] · 강민정^{††}
 Inhwan Yoo[†] · Minjeong Kang^{††}

요약

본 연구는 생성형 AI가 학습자의 사고 과정과 정보교육의 핵심 역량을 어떻게 변화시키고 있는가라는 문제의식에서 출발하였다. 생성형 AI 환경에서 학습자는 단순히 정보를 수집하거나 산출물을 생성하는 데 그치지 않고, AI 산출물을 검토하고, 판단 근거를 스스로 점검해야 한다. 이에 본 연구는 정보교육의 핵심 역량을 컴퓨팅 사고력(CT)의 재해석과 확장의 관점에서 검토하고, 인간-AI 상호작용에 요구되는 역량으로 'AI 맥락 다관점 메타인지(AI-SPM)'를 제안하였다. 선행연구 분석 결과 CT, 4C, AI 리터러시, 자기조절학습, 메타인지는 생성형 AI 환경에서의 인간-AI 상호작용의 역동성을 설명하는 데 한계가 있었다. 본 연구는 CT를 인식론적으로 재해석하여 AI 산출물의 타당성을 검토하는 기반으로 삼고, 이를 상황인지, 분산인지 및 관점주의와 결합하여 상호작용 전반에서 작동하는 메타인지적 조절 체계인 AI-SPM을 개념화하였다. AI-SPM은 기존의 자기조절학습이나 AI 리터러시와 달리, 개별 AI 산출물의 구조와 절차, 산출물과 질문에 내재된 관점, 그리고 인간-AI 상호작용 전반에 걸친 인지적 위탁 수준을 통합적 조절 대상으로 삼는다는 점에서 차별화된다. 이를 통해 정보교육의 목표를 AI 활용 능력에서 'AI와 함께 사고하는 과정을 점검하고 조절하는 능력'으로 확장하고, 과정 중심 수업 설계와 판단 근거 중심 평가의 필요성을 제시하였다. 다만 본 연구는 문헌 및 이론적 논의에 기반한 개념 제안이라는 한계를 지니며, 향후 측정 도구 개발과 실증 연구를 통한 정교화가 요구된다.

주제어 생성형 AI, 컴퓨팅사고력, 인식론적 CT, 인간-AI 상호작용, AI 맥락 다관점 메타인지(AI-SPM)

ABSTRACT

This study examines how generative AI is transforming learners' thinking processes and the discourse on core competencies in computing and information education. In generative AI environments, learners must move beyond collecting information or producing outputs to examining AI-generated outputs and evaluating the grounds for their judgments. Accordingly, this study reinterprets and extends Computational Thinking (CT) and proposes AI-Situated Perspectival Metacognition (AI-SPM) as a competency required for human-AI interaction. A review of prior research indicates that CT, the 4Cs, AI literacy, self-regulated learning, and metacognition have limitations in explaining the dynamic nature of such interaction. This study reinterprets CT from an epistemological perspective as a basis for evaluating AI-generated outputs and integrates it with situated cognition, distributed cognition, and perspectivism to conceptualize AI-SPM as a metacognitive regulatory system operating throughout human-AI interaction. Unlike self-regulated learning or AI literacy, AI-SPM regulates three dimensions: the structure and procedures underlying AI-generated outputs, the perspectives embedded in outputs and prompts, and the degree of cognitive delegation to AI. It expands the goal of computing and information education from using AI to monitoring and regulating the process of thinking with AI, and highlights the need for process-oriented instructional design and assessment centered on the grounds for learners' judgments. As a literature-based theoretical proposal, AI-SPM requires further refinement through measurement development and empirical research.

Keywords Generative AI, Computational Thinking, Epistemological CT, Human-AI Interaction, AI-Situated Perspectival Metacognition (AI-SPM)

†중신회원 대구교육대학교 컴퓨터교육과 교수
 ††정회원 대구교육대학교 교육전문대학원 AI교육
 전공 박사과정(교신저자)

논문투고 2026년 06월 03일
 심사완료 2026년 07월 20일
 게재확정 2026년 07월 23일
 발행일자 2026년 08월 27일

1. 서론

기존 정보교육의 핵심 역량 논의는 컴퓨팅사고력(Computational Thinking, 이하 CT)을 중심으로 전개되었으며, 이후 CT는 K-12 교육과 범교과적 문제 해결 역량 논의로 확장되었다[1-4]. 이러한 흐름과 함께 4C 역량은 창의성, 비판적 사고, 의사소통, 협업을 핵심 요소로 강조하였으며, AI 리터러시는 AI를 이해하고 활용하며 평가하는 능력을 중심으로 제시되어 왔다[5-7]. 이들 논의는 정보교육의 핵심 역량을 확장하고 AI 시대의 새로운 교육 요구를 반영하였다는 점에서 중요한 의의가 있다. 그러나 생성형 AI 환경에서는 학습자가 질문을 조정하고 AI 산출물의 전제와 관점을 비교하며 자신의 판단을 수정해 나가는 상호작용 과정이 중요해지고 있음에도, 기존 역량 프레임워크는 이러한 역동적 사고 과정을 충분히 설명하지 못한다.

생성형 AI의 확산은 학습자가 지식을 탐색하고 가공하며 검증하는 방식을 변화시키고 있다. 학습자는 더 이상 검색 결과를 수집하는 데 그치지 않고, AI에 문제 상황을 설명하고 산출물을 바탕으로 질문과 판단을 수정하는 과정에 참여한다. 이 과정에서 학습자는 AI와의 상호작용을 조율하는 판단 주체로 위치하게 된다[8, 9]. 따라서 생성형 AI 환경에서 요구되는 역량은 결과물을 만들어 내는 능력을 넘어, 질문, 응답, 판단의 관계를 점검하고 조절하는 능력으로 확장된다.

본 연구의 목적은 생성형 AI 시대 정보교육의 핵심 역량을 인간-AI 상호작용에서의 판단과 조절이라는 관점에서 재구성하는 데 있다. 이를 위해 첫째, CT, 4C, AI 리터러시, 자기조절학습 등 기존 핵심 역량 논의와 메타인지 논의가 생성형 AI 환경에서 갖는 의의와 한계를 검토한다. 둘째, 상황인지와 분산인지, 관점주의, 메타인지 이론을 바탕으로 인간-AI 상호작용에서 나타나는 판단과 조절의 성격을 분석한다. 셋째, 이러한 논의를 바탕으로 생성형 AI 환경에서 요구되는 학습자 역량을 새롭게 개념화한다.

2. 이론적 배경

2.1 상황인지와 분산인지

초기 인지과학은 학습을 개인 내부의 정보처리 과정으로 간주했으나, 1980년대 후반부터 인지를 실제 활동과 맥락과 연계하여 파악하려는 상황인지 관점이 대두되었다. Brown 외[10]는 학습이 활동, 맥락, 문화와 분리될 수 없음을 강조하며 탈맥락적 지식의 한계를 비판하였다. 지식은 실제 활용되는 맥락 및 활동과 결합할 때 비로소 유의미한 도구로 기능한다는 것이다.

Lave와 Wenger[11]는 이를 사회적 차원으로 확장하여 학습을 공동체 내의 참여 과정으로 재정의하였다. 학습자는 정당한 참여를 통해 점차 중심 구성원으로 이동하며, 이 과정에서 고립된 개인이 아닌 도구, 동료, 환경과 상호작용하

는 인지적 주체로 규정된다.

이러한 상황인지 이론은 Hutchins[12]의 분산인지 이론에 이르러 인지의 단위를 시스템 전체로 확장한다. 그는 해군 함정의 항법(Navigation) 업무 분석을 통해 배를 정박하는 과정이 한 명의 천재적인 항해사에 의해 이루어지는 것이 아니라, 여러 선원과 계측 도구 간의 협업을 통해 계산되고 수행됨을 규명하였다. 이후 Hollan 외[13]는 이 관점을 인간-컴퓨터 상호작용 연구에 적용하여, 인간과 디지털 도구가 형성하는 통합적인 인지 시스템 자체를 분석 단위로 삼아야 한다고 주장하였다.

종합하면, 상황인지와 분산인지 관점은 인지를 활동과 도구, 타 행위자 간의 상호작용 속에서 구성되는 과정으로 이해한다. 이러한 관점은 생성형 AI 산출물이 인간의 사고와 분리된 결과물이 아니라, 인간-AI 상호작용 과정에서 형성되는 인지적 산출물임을 이해하는 데 중요한 시사점을 제공한다.

2.2 관점주의

관점주의는 인간의 인식과 이해가 특정한 관점을 통해 이루어진다는 점을 강조한다. Nietzsche[14]는 인간이 세계를 있는 그대로 인식할 수 있다는 생각에 의문을 제기하며, 모든 이해는 주체의 시각과 가치에 의해 형성된다고 보았다. Haraway[15] 역시 지식은 특정한 위치와 맥락 속에서 형성된다고 설명하며, 완전히 중립적이고 객관적인 시선의 가능성을 비판하였다.

그러나 관점주의가 모든 관점을 동일하게 인정하는 상대주의를 의미하는 것은 아니다. Hales와 Welshon[16]은 관점주의가 단일한 시각의 한계를 인식하고, 서로 다른 관점을 비교하고 조정하는 과정을 통해 보다 타당한 이해에 접근하려는 인식적 태도를 요구한다고 설명하였다. 즉, 관점주의의 핵심은 다양한 관점의 존재를 인정하는 데 그치지 않고, 이를 비판적으로 검토하고 조정하는 과정에 있다.

이러한 관점주의의 입장은 생성형 AI 산출물을 이해하는 데 중요한 시사점을 제공한다. 생성형 AI의 산출물은 모델이 세계를 이해한 결과라기보다, 학습 데이터와 설계 방식, 조정 과정의 영향을 받아 생성된 결과이다[17]. 그럼에도 생성형 AI의 유창성으로 인해 학습자가 이를 객관적 사실이나 중립적 판단으로 받아들일 가능성이 존재한다. 관점주의는 이러한 위험성을 경계하며, 학습자가 AI 산출물에 포함된 전제와 관점을 비판적으로 검토하고 비교할 필요가 있음을 시사한다.

3. 핵심 역량에 대한 선행연구 분석

3.1 기존 역량 프레임워크

3.1.1 컴퓨팅사고력(CT)

컴퓨팅사고력(Computational Thinking, 이하 CT)은

정보 및 컴퓨터 교육 분야에서 핵심적으로 다루어져 온 대표적 역량이다. Wing[1]은 CT를 컴퓨터과학의 기본 개념을 활용하여 문제 해결, 시스템 설계, 인간 행동 이해를 수행하는 사고방식으로 제시하였다. 이후 교육 분야에서 CT는 문제를 분해하고, 패턴을 인식하며, 핵심 요소를 추상화하고, 해결 절차를 알고리즘적으로 설계하는 사고 과정으로 구체화되었다[2-4].

CT는 단순한 프로그래밍 기술이 아니라, 컴퓨팅 기반 문제 해결 방식을 교육 가능한 형태로 구조화한 사고 체계라는 점에서 중요한 의미를 지닌다. CT는 문제를 계산 가능한 형태로 분석하고, 해결 과정을 명시적인 절차로 표현하도록 돕는다. 이와 함께 CT는 컴퓨팅 시스템의 처리 가능성과 한계를 이해하고, 문제 해결 과정에서 인간의 판단과 컴퓨터의 수행이 어떻게 분담되는지를 파악하는 사고와도 관련된다.

이러한 특성으로 인해 CT는 프로그래밍 교육을 넘어 수학, 과학, 공학 등 다양한 교과와 결합하며 범교과적 문제 해결 역량으로 확장되었다[2-4]. 한편 Brennan과 Resnick[18]은 CT를 개념, 실천, 관점의 차원으로 설명하며 산출물 구조 이해와 오류 수정, 설계 과정 성찰을 강조하였다. 이는 CT가 단순한 문제 해결 절차의 구성을 넘어 산출물과 설계 과정을 점검하는 성찰적 사고의 틀로 확장되어 왔음을 보여준다. 그러나 생성형 AI 환경에서는 문제 해결 과정이 명시적 절차의 설계를 넘어, AI 응답의 타당성을 판단하고 그 활용 범위를 조정하는 상호작용으로 확장된다는 점에서 기존 CT 논의만으로는 설명되지 않는 영역이 남는다.

3.1.2 4C 역량

4C(창의성, 비판적 사고, 의사소통, 협업)는 21세기 핵심 역량 논의를 주도해 온 대표적인 프레임워크이다[5]. 창의성과 비판적 사고는 새로운 아이디어를 산출하고 정보와 근거의 타당성을 판단하는 역량과 관련되며, 의사소통과 협업은 공동의 목표 달성을 위해 의미를 공유하고 타인과 상호작용하는 역량을 포함한다.

4C는 정보교육에서도 프로그래밍, 설계, 프로젝트 기반 활동과 결합하여 학습자의 창의적 산출과 협력적 문제 해결을 해석하는 일반 역량 틀로 활용될 수 있다[9]. 이러한 가치는 생성형 AI 환경에서도 여전히 유효하며, AI를 활용한 창의적 산출, AI 응답에 대한 비판적 검토, 인간과 AI의 상호작용을 통한 의사소통과 협업은 중요한 교육 과제로 부상하고 있다[8]. 그러나 4C는 범용적 역량 범주를 중심으로 구성되어 있어, 생성형 AI와의 상호작용에서 학습자가 어떤 기준으로 AI 산출물을 검토하고 자신의 판단을 조정해야 하는지를 구체적으로 설명하는 데에는 한계가 있다. 특히 AI 산출물에 포함된 다양한 관점을 비교하고, 인간의 판단과 AI의 제안을 조율하는 과정은 4C의 일반적 틀만으로 체계화하기 어렵다.

3.1.3 AI 리터러시

AI 리터러시는 최근 교육 논의에서 핵심 역량으로 부상

하고 있으며, AI를 이해하고 책임 있게 활용하는 데 요구되는 지식, 기술, 태도를 포괄한다. Allen과 Kendeou[6]은 이를 개념적, 절차적, 비판적, 윤리적 이해로 구성된 다차원 역량으로 제시하였다. UNESCO[7]의 주요 프레임워크 역시 AI 시스템의 원리, 데이터와 알고리즘의 편향성, 사회적 영향 등을 핵심 요소로 다루며, AI 산출물의 신뢰성과 오용 가능성에 대한 비판적 인식을 강조한다.

이러한 논의는 AI를 단순한 활용 도구가 아니라 비판적 성찰의 대상으로 확장하였다는 점에서 중요한 의미를 지닌다. 그러나 대부분의 AI 리터러시 프레임워크는 지식 이해와 태도 형성 수준에 초점을 두고 있어, 실제 인간-AI 상호작용 과정에서 학습자가 자신의 판단을 어떻게 점검하고 조정해야 하는지에 대한 설명은 상대적으로 부족하다. 예를 들어 AI 산출물을 비판적으로 평가해야 한다는 원칙은 제시되지만, 해당 평가 과정이 어떠한 메타인지 전략을 통해 이루어져야 하는지에 대한 구체적 틀은 충분히 제시되지 않고 있다.

3.1.4 자기조절학습

자기조절학습(Self-Regulated Learning, SRL)은 학습자가 목표를 설정하고 전략을 실행하며 수행 결과를 성찰함으로써 학습을 주도적으로 이끄는 과정을 말한다. Zimmerman[19]은 자기조절학습을 사전계획, 수행, 자기성찰의 순환 구조로 제시하였으며, 학습자가 자신의 인지적, 동기적, 행동적 과정을 조정하면서 학습을 개선한다고 보았다. 이는 메타인지가 특정 지식을 익히는 상황에 국한되지 않고, 목표 설정부터 최종 성찰까지 학습 전반에서 작동함을 보여준다.

다만 자기조절학습은 주로 학습자 개인의 내적 조절을 중심으로 다루어졌다. 생성형 AI 환경에서는 외부 생성 도구가 제공하는 산출물과 이를 매개로 한 상호작용이 학습자의 사고 형성에 직접 개입할 수 있으므로, 조절의 범위도 학습자 내부 과정에 한정되기 어렵다. 따라서 자기조절학습은 생성형 AI 환경에서 요구되는 순환적 조절 구조를 이해하는데 유용한 기반을 제공하지만, 조절 대상이 AI 산출물과 인간-AI 상호작용 과정으로 확장될 필요가 있다.

3.1.5 메타인지

메타인지는 학습자가 자신의 인지 과정을 인식하고 조절하는 능력을 가리킨다. Flavell[20]은 메타인지를 인지 현상에 대한 지식과 인식으로 정의하며, 메타인지적 지식, 메타인지적 경험, 목표와 과제, 행동과 전략 간의 상호작용을 설명하였다. Brown[21]은 메타인지를 인지에 대한 지식과 인지 조절의 두 측면으로 구분하였으며, 계획, 점검, 평가와 같은 조절 과정을 핵심 요소로 제시하였다. 또한 Nelson과 Narens[22]는 실제 인지 수행과 이를 점검하고 통제하는 상위 인지 간의 순환적 상호작용 구조를 제안함으로써 메타인지의 조절 과정을 체계화하였다.

이와 같은 고전적 메타인지 이론은 학습자가 자신의 인지 과정을 점검하고 조절하는 구조를 설명하는 데 중요한 토대를 제공하였다. 그러나 후속 연구들은 메타인지가 탈맥락적이고 보편적인 상위 인지 능력에 머무르지 않으며, 특정 교과 내용, 도구, 과제, 상호작용 속에서 발현되는 상황적 역량임을 강조해 왔다[23]. 따라서 메타인지는 개인 내부의 사고 점검을 넘어, AI 산출물의 전제와 관점, 자신의 판단 기준, 맥락에 따른 수용 가능성을 함께 점검하고 조율하는 과정으로 확장될 필요가 있다.

3.2 선행연구 고찰의 시사점

CT, 4C, AI 리터러시, 자기조절학습, 메타인지에 대한 논의는 정보교육에서 학습자의 핵심 역량을 설명하는 데 중요한 기반을 제공해 왔다. 그러나 생성형 AI가 학습 환경 전반에 확산됨에 따라, 기존 논의가 전제해 온 학습자, 도구, 맥락의 관계는 재검토될 필요가 있다. 특히 기존 역량 논의의 한계는 다음 두 측면에서 드러난다.

첫째, 상황성에 대한 설명이 제한적이다. 생성형 AI 환경에서 학습자는 AI에 문제 상황을 설명하고, 응답을 해석하며, 동료와 의미를 조율하고, 과제 맥락에 맞게 산출물을 수정한다. 이 과정에서 사고는 개인 내부에서만 이루어지는 것이 아니라, 학습자, AI, 동료, 과제, 디지털 도구가 결합된 상호작용 속에서 형성된다[10, 12, 13]. 따라서 생성형 AI 환경의 핵심 역량을 논의하기 위해서는 개별 학습자의 능력뿐 아니라 인간-AI-과제-환경이 결합된 상호작용 체계를 함께 고려할 필요가 있다.

둘째, 관점성에 대한 설명이 충분히 구조화되어 있지 않다. 생성형 AI 산출물은 객관적이고 중립적인 결과처럼 제시되지만, 앞서 살펴본 것처럼 데이터와 설계의 영향 속에서 구성된 산출물이다[17]. 따라서 학습자는 AI 산출물의 오류 여부만이 아니라, 그 안에 포함된 전제와 관점, 누락된 입장, 자신의 판단 기준을 함께 검토해야 한다. 그러나 기존 역량 논의는 이러한 관점 비교와 판단 재구성의 과정을 체계적으로 설명하는 데 한계를 보인다.

이처럼 생성형 AI 환경은 학습자에게 맥락을 파악하고, 다양한 관점을 비교하며, 인간의 판단과 AI의 제안을 조율하는 사고를 요구한다. 따라서 AI 시대의 정보교육은 기존 역량을 병렬적으로 추가하는 방식이 아니라, 상황성, 관점성, 조절의 문제를 통합적으로 설명하는 방향으로 재구성될 필요가 있다.

4. AI 시대 CT와 메타인지에 대한 재해석

4.1 CT의 재해석에 대한 논의

생성형 AI의 등장은 AI 산출물의 생성 논리와 타당성을 검토하는 사고로서 CT의 역할을 재해석할 필요를 제기한다. 이에 본 연구는 CT를 실행적 CT와 인식론적 CT로 구분하여 설명하고자 한다. 실행적 CT는 학습자가 문제를 분해

하고, 패턴을 인식하며, 핵심 요소를 추상화하고, 해결 절차를 구성하는 등 문제 해결 과정을 직접 수행하는 CT의 작동 양상을 의미한다. 반면 인식론적 CT는 AI가 제시한 산출물의 구조, 근거, 추상화 수준, 절차적 타당성을 CT의 기준으로 검토하는 작동 양상을 의미한다. 즉, 인식론적 CT는 CT를 산출물 생성의 절차가 아니라, AI 산출물의 생성 논리와 타당성을 검토하는 판단 기준으로 재위치화한 개념이다.

이 구분의 의의는 기존 CT 및 AI 리터러시 논의와의 관계 속에서 분명해진다. 기존 논의는 AI 시스템의 신뢰성, 편향성, 윤리성, 사회적 영향에 대한 비판적 이해와 책임 있는 활용을 강조하지만[6, 7], 개별 AI 산출물의 분해 구조, 추상화 수준, 절차적 타당성을 CT의 언어로 분석하는 사고 절차를 충분히 구체화하지는 않는다. 따라서 인식론적 CT는 기존 CT와 AI 리터러시를 대체하는 개념이라기보다, 생성형 AI 산출물 판단을 위해 CT의 핵심 요소를 재배치한 분석 틀로 이해할 수 있다. 인식론적 CT와 기존 CT 및 AI 리터러시 논의의 관계는 Table 1과 같다.

Table 1. Relationship between Epistemological CT, Existing CT, and AI Literacy Discourse

Category	Object of Judgment	Criteria for Judgment
Brennan and Resnick's CT Perspectives Dimension	Design processes and artifacts directly produced by learners	Alignment between design intent and artifact; reflection on the process
Critical Evaluation in AI Literacy	Reliability, bias, and ethics of AI systems and outputs	System understanding; data and algorithmic bias; social impact and accountability
Epistemological CT	Structure, procedure, and reasoning of AI-generated outputs	CT criteria: decomposition, pattern recognition, abstraction, and algorithm design

즉 실행적 CT와 인식론적 CT는 새로운 CT 구성 요소의 추가가 아니라, 기존 CT 요소가 생성형 AI 환경에서 산출물 생성과 산출물 판단이라는 서로 다른 방향으로 작동하는 양상을 설명하기 위한 구분이다. 실행적 CT에서 분해, 패턴 인식, 추상화, 알고리즘 설계는 학습자가 문제 해결 과정을 직접 구성하는 데 사용된다. 반면 인식론적 CT에서 동일한 요소들은 AI 산출물이 문제를 어떻게 구조화했는지, 어떤 오류 패턴을 보이는지, 어느 수준의 추상화에 머물러 있는지, 해결 절차가 타당한지를 검토하는 판단 기준으로 사용된다. 이를 CT의 주요 요소에 따라 정리하면 Table 2와 같다.

Table 2. Comparison between Executional CT and Epistemological CT

Core CT Element	Executional CT	Epistemological CT
Decomposition	Breaks a problem into sub-components and identifies their relationships	Evaluates whether an AI output appropriately decomposes a problem into core and peripheral conditions

Core CT Element	Executorial CT	Epistemological CT
Pattern Recognition	Identifies recurring patterns and generalizes them	Detects recurring errors, biases, and omissions
Abstraction	Extracts key elements and omits unnecessary details	Judges whether the level of generalization is appropriate to the task context
Algorithm Design	Constructs solution procedures or algorithms	Examines the procedural validity of the solution process and reasoning chain presented by AI

인식론적 CT에서 분해는 AI 산출물이 문제의 핵심 조건과 부수 조건을 적절히 구분하고 있는지를 평가하는 기준이 된다. 패턴 인식은 반복적으로 나타나는 오류, 편향, 누락, 과잉 일반화의 양상을 파악하는 데 활용된다. 추상화는 AI 산출물이 과제 맥락에 비추어 지나치게 일반적이거나 세부적인 수준에 머물러 있지 않은지를 판단하게 한다. 알고리즘 설계는 AI가 제시한 해결 과정, 논리 흐름, 근거 연결이 절차적으로 타당하게 구성되어 있는지를 검토하는 기준이 된다.

인식론적 CT는 Brennan과 Resnick[18]의 CT ‘관점’차원과 구별될 필요가 있다. 두 개념은 모두 산출물과 과정에 대한 성찰을 포함한다는 점에서 연결되지만, 성찰의 대상이 다르다. Brennan과 Resnick의 관점 차원에서 학습자가 성찰하는 대상은 자신이 직접 설계한 코드와 프로젝트 산출물이며, 성찰의 초점은 설계 의도와 산출물 간의 정합성에 있다. 반면 인식론적 CT에서 판단 대상은 생성형 AI가 산출한 결과물로 전환되며, 학습자는 자신이 만들지 않은 산출물의 분해 구조, 추상화 수준, 절차적 타당성을 CT의 기준으로 역분석한다. 따라서 인식론적 CT는 관점 차원의 ‘성찰’ 대상을 자기 산출물에서 타자(AI)의 산출물로 확장하고, 성찰의 준거를 설계 의도 정합성에서 산출물의 구조적, 절차적 타당성으로 재설정한다는 점에서 차별화된다.

인식론적 CT는 생성형 AI 환경에서 더욱 중요해진다. AI 산출물은 유창하고 완성도 높은 형태로 제시될 수 있지만, 그것이 곧 정확성이나 적절성을 보장하는 것은 아니다. 따라서 이를 실제 학습 맥락에서 언제, 어떤 방식으로 적용하고 조율할 것인가는 별도의 문제로 남는다. 이러한 점에서 인식론적 CT는 과제 맥락과 AI 산출물의 관점을 함께 고려하며 학습자의 판단 기준을 조율하는 상위 수준의 메타인지 역량과 결합될 때 교육적 의미가 더욱 분명해진다.

4.2 생성형 AI 환경에서 메타인지 개념의 확장

전통적인 학습 환경에서 메타인지는 주로 학습자가 수행 결과를 점검하고 조절하는 과정으로 설명되어 왔으나, 생성형 AI 활용 과정에서 메타인지는 학습자 자신의 이해뿐 아니라 AI와의 상호작용이 판단을 어떻게 형성하고 변화시키는지를 점검하는 방향으로 확장된다.

따라서 생성형 AI 환경에서 메타인지는 학습자 개인의 사고 과정에 대한 점검을 넘어, 인간-AI 상호작용의 맥락,

산출물의 관점성, AI에 대한 의존 수준을 조절하는 능력으로 확장될 필요가 있다.

생성형 AI 활용은 특정한 과제 맥락과 분리되어 판단될 수 없다. 상황인지와 분산인지 관점에서 볼 때 인지는 활동, 도구, 맥락 속에서 형성되며[10, 12, 13], 생성형 AI 활용 역시 과제 목적, 수업 목표, 평가 기준, 도구 환경 속에서 이루어진다. 따라서 같은 AI 활용이라도 목적이 개념 이해인지, 아이디어 생성인지, 결과물 개선인지에 따라 적절한 활용 범위가 달라질 수 있으며, 학습자는 과제의 목적과 맥락에 비추어 AI 활용 방식의 적절성을 검토해야 한다.

또한 학습자는 AI 산출물에 내재된 관점뿐 아니라, 자신의 질문에 전제된 관점까지 성찰할 수 있어야 한다. 인식은 특정한 위치와 맥락 속에서 형성되며[14, 15], 생성형 AI 산출물 역시 학습 데이터와 설계 방식의 영향을 받아 구성된 결과이다[17]. 따라서 학습자는 이를 비판적으로 검토하며 능동적으로 대안을 탐색할 필요가 있다.

이러한 점검은 AI에 대한 의존 수준을 조절하는 문제로 이어진다. 생성형 AI는 학습자의 인지적 부담을 낮출 수 있지만, 학습자가 문제 해결의 구조와 판단 근거를 충분히 이해하지 못한 채 결과를 수용하게 할 위험도 있다[24, 25]. 따라서 학습자는 어떤 과정을 AI에 위탁하고 어떤 판단은 스스로 유지해야 하는지를 조절할 수 있어야 하며, 이러한 조절이 실제 판단으로 이어지기 위해서 AI 산출물의 구조와 조건을 파악하는 사고가 필요하다.

4.3 인식론적 CT와 확장된 메타인지의 통합

생성형 AI 환경에서 학습자의 판단은 AI 산출물을 분석하는 문제와 인간-AI 상호작용을 조절하는 문제를 동시에 포함한다. 따라서 인식론적 CT와 확장된 메타인지의 통합은 두 개념을 병렬적으로 결합하는 것이 아니라, 산출물 분석의 기준과 상호작용 조절의 조건이 어떻게 연결되는지를 밝히는 것이다.

인식론적 CT는 확장된 메타인지가 구체적인 판단으로 작동하기 위한 분석 기준을 제공한다. 학습자가 AI 산출물을 점검한다고 하더라도, 무엇을 기준으로 오류, 누락, 부적절한 단순화, 절차적 비약을 판단할 수 있는지가 분명하지 않으면 점검은 일반적 경계나 주관적 인상에 머무를 수 있다. 이때 문제의 하위 요소, 핵심 조건, 절차와 결과의 관계를 파악하는 CT는 AI 산출물이 과제 요구를 충족하는지, 어떤 과정이 생략되거나 축약되었는지, 수정 또는 재질문이 필요한 지점이 어디인지를 판단하게 한다. 생성형 AI는 컴퓨팅 교육에서 코드 생성, 설명 제공, 오류 수정 지원 등 학습 산출 과정의 일부에 직접 개입하고 있다[9]. 이러한 상황에서 인식론적 CT는 학습자가 AI에 의해 대신 수행된 문제 분해, 절차 구성, 근거 연결을 검토하고, 필요한 지점에서 인간의 판단을 개입시키는 기준으로 기능한다.

반대로 확장된 메타인지는 인식론적 CT가 개별 산출물의 구조 분석에만 머무르지 않고 인간-AI 상호작용의 조절로 이어지도록 한다. AI 산출물이 구조적으로 그럴듯하더라도

도 그것이 모든 과제 맥락에서 적절한 것은 아니며, 동일한 분석 기준도 개념 이해, 아이디어 생성, 결과물 개선과 같은 목적에 따라 다르게 적용될 수 있다. 또한 산출물의 전제와 관점, 학습자의 질문 방식, AI 제안에 대한 의존 수준은 CT 분석의 범위와 강도를 조정한다. 따라서 확장된 메타인지는 인식론적 CT가 언제, 어떤 방향으로, 어느 정도 적용되어야 하는지를 결정하는 적용 조건을 제공한다.

이와 같은 맥락에서 생성형 AI 시대의 새로운 역량 개념으로서 본 연구는 이를 AI 맥락 다관점 메타인지(AI-Situated Perspectival Metacognition, AI-SPM)로 제안하며, 다음 장에서 그 개념과 구성 요소를 구체화한다.

5. AI-SPM의 개념과 구성 요소

5.1 AI-SPM의 개념

AI 맥락 다관점 메타인지(AI-SPM)는 생성형 AI와의 상호작용 속에서 학습자가 자신의 사고, AI 산출물, 그리고 인간-AI 상호작용 과정을 맥락과 관점의 측면에서 점검하고 조절하는 상위 메타인지 역량이다.

AI-SPM에서 ‘맥락(Situated)’은 AI 활용이 과제 목적, 평가 기준, 수업 규범, 도구 환경에 따라 달라진다는 점을 의미하며, 학습 상황에서 AI 활용 방식과 위탁 범위가 적절한지를 점검하고 조절하는 역량을 포함한다.

또한, 학습자는 AI 산출물의 내재적 관점을 점검할 필요가 있으며, 학습자 스스로 본인의 질문 관점을 성찰해야 한다. 특히 산출물을 그대로 수용하기보다 능동적으로 새로운 관점으로 문제에 접근하는 자세가 필요하다. 따라서 ‘다관점(Perspectival)’은 첫째, AI 산출물의 내재적 관점 점검, 둘째, 학습자 질문의 관점 성찰, 셋째, 대안적 관점의 능동적 생성이라는 세 가지 범주로 구분한다.

AI-SPM에서 ‘메타인지(Metacognition)’는 인간-AI 상호작용 과정 전체를 대상으로 하는 점검과 조절을 의미한다. 학습자는 AI 산출물을 단순히 수용하거나 수정하는 것이 아니라, 그 산출물이 자신의 사고에 어떤 영향을 미쳤는지, 어떤 판단이 자신의 이해에 기반한 것인지, 어떤 판단이 AI의 제안에 의해 형성된 것인지 성찰해야 한다.

AI-SPM은 CT를 대체하는 것이 아니라, 인식론적 CT가 제공하는 산출물 판단의 기반을 인간-AI 상호작용 전체의 메타인지적 조절로 확장한 개념이다. 다음 절에서는 AI-SPM이 어떠한 이론적 흐름을 바탕으로 도출되는지 살펴보고, 그 구성 요소를 구체화하고자 한다.

5.2 AI-SPM의 도출과 구성 요소

생성형 AI 환경에서 학습자는 AI 산출물을 단순히 사용하는 수행자에 머무르지 않고, AI와 상호작용하며 산출물의 의미와 한계를 해석하는 조정자로 변화한다. 본 연구는 이러한 변화를 바탕으로 CT 및 메타인지의 대상을 확장하여 AI-SPM의 개념적 기반을 도출하였다. AI-SPM의 도출 논리와

이론적 흐름은 Fig. 1과 같다(상세한 내용은 부록#1 참조).

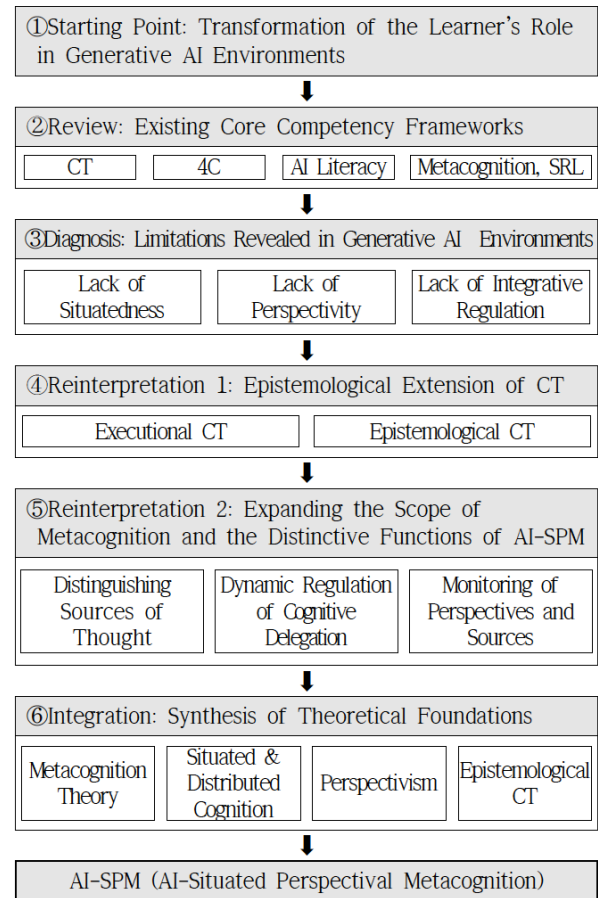


Figure 1. Conceptual Derivation of AI-SPM

AI-SPM은 네 가지 이론적 기반을 통해 도출된다. 첫째, 메타인지 이론은 지식, 감시, 조절이라는 기능 축을 제공한다. 둘째, 상황인지와 분산인지는 인간-AI 상호작용을 상호작용 전, 중, 후, 장기 축적의 과정으로 이해하는 기반을 제공한다. 셋째, 관점주의는 AI 산출물과 학습자의 질문에 포함된 전제, 관점, 누락에 대한 점검뿐 아니라 능동적으로 대안을 탐색할 수 있는 근거를 제공한다. 넷째, 인식론적 CT는 AI 산출물의 구조, 근거, 추상화 수준, 절차적 타당성을 판단하는 기준을 제공한다. 이를 정리하면 Table 3과 같다.

Table 3. Theoretical Foundations and Derivational Axes of AI-SPM

Theoretical Foundation	Contribution to the Derivation of AI-SPM Components
Metacognition Theory	Provides the functional structure of knowledge, monitoring, and regulation
Situated & Distributed Cognition	Provides the procedural structure across pre-, in-, post-, and long-term accumulation phases of human-AI interaction
Perspectivism	Provides the grounds for examining assumptions, perspectives, and omissions in outputs and prompts

Theoretical Foundation	Contribution to the Derivation of AI-SPM Components
Epistemological CT	Provides the criteria for judging the structure, grounds, abstraction level, and procedural validity of AI outputs

이론적 기반의 네 개의 축은 AI-SPM의 각 구성 요소 안에서 중첩적으로 작동한다. 예컨대 하나의 구성 요소는 메타인지의 지식, 감시, 조절 중 하나에 해당하면서도, 동시에 인간-AI 상호작용의 특정 국면에서 작동하고, 산출물의 관점이나 구조적 타당성을 점검하는 역할을 수행할 수 있다. 따라서 본 연구는 이러한 이론적 기반을 바탕으로 AI-SPM의 구성 요소를 상호작용 전, 상호작용 중, 상호작용 후, 장기 축적의 네 국면에 따라 정리하였다. 이를 제시하면 Table 4와 같다.

Table 4. Components of AI-SPM

	Knowledge	Monitoring	Regulation
Pre-interaction	①Epistemological Knowledge	②Monitoring of Output Assumptions	③Planning of AI Interaction
In-interaction	④Knowledge of Human CT Competency	⑤Monitoring of Structural Validity	⑥Regulation of Prompt Reconstruction
Post-interaction		⑦Evaluation of Human-AI Interaction	⑧Regulation of AI Cognitive Delegation Scope ⑨Active Generation of New Perspectives
Long-term Accumulation	⑩Monitoring of Integrated Development of CT-AI Competencies		

상호작용 전 국면은 AI 활용을 시작하기 전에 필요한 준비와 계획 기능을 포함한다. ①인식론적 지식은 학습자가 AI 산출물의 생성 방식, 한계, 오류 가능성을 이해하고 산출물을 무비판적으로 수용하지 않도록 하는 사전 지식이다. ②산출물 전제 감시는 산출물이 제시된 이후의 검토만이 아니라, 자신의 질문이 어떤 산출물을 유도할 수 있는지, 그 질문이 어떤 전제와 관점을 포함하는지를 사전에 점검하는 기능을 포함한다. ③AI 상호작용 계획은 과제 목적과 학습 목표에 따라 AI 활용 범위와 상호작용 방식을 미리 설정하는 조절 기능이다.

상호작용 중 국면은 AI 산출물을 검토하고, 그 결과에 따라 상호작용 방식을 조정하는 기능을 포함한다. ④인간-CT 역량 지식은 학습자가 자신의 CT 역량으로 AI 산출물을 어느 수준까지 판단할 수 있는지를 인식하는 지식이다. ⑤구조적 타당성 감시는 AI 산출물의 구조, 근거, 추상화 수준, 논리 흐름이 과제 맥락에 적절한지를 검토하는 기능이다. 이때 점검 대상은 학습자 자신의 수행이 아니라 외부 행위자인 AI가 생성한 산출물이며, 점검 기준 또한 일반적 목표 부합 여부가 아니라 인식론적 CT가 제공하는 구조, 근거, 추

상화 수준, 절차적 타당성이라는 점에서 자기조절학습의 감시와 구별된다. ⑥프롬프트 재구성 조절은 이러한 감시 결과를 바탕으로 질문의 방향, 조건, 전제, 추상화 수준을 조정하는 기능이다.

상호작용 후 국면은 AI와의 상호작용 결과와 과정을 되돌아보고, 이후의 AI 활용 방식을 조정하는 기능을 포함한다. ⑦인간-AI 상호작용 평가는 하나의 상호작용이 끝난 뒤 과정과 결과가 적절했는지를 되돌아보는 감시 기능이다. 이때 AI의 응답이 자신의 문제 이해, 사고 방향, 판단 기준에 어떤 영향을 미쳤는지도 함께 평가된다. ⑧AI 인지 위탁 범위 조절은 AI에게 맡길 수 있는 인지 과정과 학습자가 직접 유지해야 할 판단을 구분하는 기능이다. ⑨새로운 관점 생성은 AI가 충분히 제시하지 못한 관점, 이해관계자 입장, 대안적 해석을 학습자가 능동적으로 생성하여 산출물의 관점적 한계를 보완하는 기능이다.

장기 축적 국면은 반복적인 AI 활용 경험이 학습자의 역량 발달에 미치는 영향을 점검하는 기능을 포함한다. ⑩CT-AI 역량 통합발달 감시는 단일 상호작용이 아니라, 여러 차례의 AI 활용 경험이 누적되는 과정에서 학습자의 CT 역량과 AI 활용 역량이 어떻게 변화하는지를 장기적으로 점검하는 기능이다. 이는 AI-SPM이 일회적 사용 전략이 아니라, AI와의 지속적 상호작용 속에서 인간의 사고 역량이 어떻게 발달하거나 약화되는지를 살피는 메타인지 체계임을 보여준다.

이처럼 AI-SPM의 구성 요소는 생성형 AI 활용의 전 과정을 따라 학습자의 판단 기준, 관점 검토, 위탁 범위 조절을 연결하는 순환적 체계로 이해할 수 있다. 기존 프레임워크가 각각 학습자 내부 인지, AI 시스템 일반, 자작 산출물 중 하나에 점점의 초점을 두었다면, AI-SPM은 이들을 병렬적으로 추가하는 것이 아니라 인간-AI 상호작용 체계를 단일한 분석 단위로 삼아 상황성, 관점성, 조절을 하나의 틀로 통합한다는 점에서 구별되며 기존 역량과의 차이점을 비교하면 Table 5와 같다.

Table 5. Differentiation between Existing Competency Frameworks and AI-SPM

Frame-work	Primary Object of Monitoring & Regulation	Unit of Analysis	Differentiation from AI-SPM
CT	Problem-solving procedures constructed by the learner	Learner-produced artifacts and code	Does not extend the object of judgment to AI-generated outputs
4Cs	Competencies of creativity, critical thinking, communication, and collaboration	General thinking and collaborative activities	Lacks concrete criteria and procedures for examining AI outputs
AI Literacy	Reliability, bias, and ethics of AI systems	AI systems in general	Lacks learner judgment and monitoring within the human-AI interaction process

Frame-work	Primary Object of Monitoring & Regulation	Unit of Analysis	Differentiation from AI-SPM
SRL	Learner's internal cognition, motivation, and behavior	The individual's performance process	Confines the object of regulation to the learner's interior, excluding AI interaction
Meta cognition	One's own cognitive processes	Decontextualized individual cognition	Does not address the perspectival nature of outputs or the blending of thought sources

5.3 AI-SPM 교육적 함의

AI-SPM은 생성형 AI 시대 정보교육의 목표가 단순한 기술 활용이나 산출물 생성 능력에 머무를 수 없음을 시사한다. 정보교육은 AI를 효과적으로 사용하는 능력을 넘어, AI와 함께 형성되는 사고 과정을 점검하고 조절하는 역량을 핵심 교육 목표로 포함할 필요가 있다.

교수-학습 설계에서는 AI 사용 자체가 아니라 AI와의 상호작용 과정이 학습의 대상이 되어야 한다. 과제는 최종 산출물 제출로 끝나기보다, 학습자가 질문의 의도, AI 산출물에 대한 검토 근거, 수정 또는 재질문이 필요한 이유를 드러내도록 구성될 필요가 있다. 이때 프롬프트를 정교하게 작성하는 기술보다 중요한 것은 학습자가 질문과 산출물의 관계를 이해하고, 그 과정에서 자신의 판단이 어떻게 조정되었는지를 설명할 수 있도록 하는 것이다. 이러한 설계는 CT를 개념, 실천, 관점이 결합된 사고 과정으로 이해하고 학습자의 성찰을 강조하는 논의[18]와도 연결된다.

평가 역시 결과물의 완성도만으로 이루어지기 어렵다. AI-SPM 관점에서 평가의 핵심 증거는 학습자가 AI 산출물을 어떤 근거로 판단했는지, 어떤 부분을 수용하거나 수정했는지, 어떤 판단을 AI에 위탁하지 않고 유지했는지에 있다. 따라서 평가는 최종 산출물뿐 아니라 상호작용 기록, 판단 근거, 인지적 위탁 범위에 대한 설명을 함께 살펴야 한다. 이는 생성형 AI 활용이 인지적 오프로딩과 자기점검 약화를 동반할 수 있다는 논의와도 관련되며, 평가의 초점이 AI 사용 여부가 아니라 학습자의 판단 유지와 조절 과정에 놓여야 함을 보여준다[24, 25].

AI-SPM은 기존 CT 중심의 정보교육에 새로운 방향을 제시한다. 기존 교육에서 문제를 직접 분해하고, 핵심 요소를 추상화하며, 해결 절차를 구성하는 활동을 강조했다면, 생성형 AI 환경에서는 AI가 생성한 산출물을 역으로 분석하고 수정하는 활동도 중요하다. 학습자는 AI 산출물이 과제 조건을 적절히 반영했는지, 추상화 수준과 절차적 연결이 타당한지를 검토하면서 실행 중심의 CT를 산출물 판단을 위한 인식론적 CT로 확장할 수 있다.

이와 같이 AI-SPM 개념은 AI 시대에 학습자가 AI와 함께 사고하는 과정에서 판단의 주도권과 책임을 유지하며 조절할 수 있도록 하는 교육의 목표, 내용, 방법, 평가의 기초

를 제공할 수 있을 것으로 기대한다.

6. 결론 및 제언

본 연구는 생성형 AI가 학습자의 사고 과정과 정보교육의 핵심 역량 논의를 어떻게 변화시키고 있는가라는 문제의식에서 출발하였다. 생성형 AI 환경에서 학습자는 더 이상 정보를 수집하거나 산출물을 생성하는 데 그치지 않고, AI에게 문제를 설명하고, AI 산출물을 검토하며, 질문을 수정하고, 최종 판단의 근거를 스스로 점검해야 한다. 이에 본 연구는 생성형 AI 시대 정보교육의 핵심 역량을 CT의 대체가 아니라 재해석과 확장의 관점에서 검토하고, 인간-AI 상호작용 속에서 요구되는 상위 메타인지 역량으로 AI 맥락 다관점 메타인지(AI-SPM)를 제안하였다.

선행연구 분석 결과, CT, 4C, AI 리터러시, 자기조절학습, 메타인지는 정보교육의 핵심 역량 논의에 중요한 기반을 제공해 왔으나, 생성형 AI 환경에서 나타나는 인간-AI 상호작용의 역동성을 충분히 설명하는 데에는 한계가 있었다. 기존 프레임워크는 대체로 학습자 개인의 역량, AI 활용 능력, 산출물의 신뢰성 평가에 초점을 두는 경향이 있다. 그러나 생성형 AI 환경에서는 AI 활용이 이루어지는 과제 맥락, AI 산출물과 학습자 질문에 포함된 전제와 관점, AI에 대한 의존 수준, 그리고 인간과 AI가 함께 형성하는 사고 과정 전체를 함께 고려해야 한다. 이는 개별 역량을 단순 누적하는 접근으로는 충분하지 않으며, 상황성, 관점성, 메타인지적 조절을 하나의 틀로 엮어내는 개념적 재구성이 필요함을 의미한다.

본 연구는 이러한 전환을 설명하기 위해 CT를 인식론적으로 재해석하였다. 인식론적 CT는 AI 산출물의 생성 구조와 추론 절차의 타당성을 검토함으로써 학습자가 AI 산출물을 무비판적으로 수용하지 않도록 돕는 판단 기반이다. 그러나 AI 산출물의 관점, 질문 방식의 영향, 사고 출처의 혼재, AI 의존 수준까지 다루기 위해서는 이를 인간-AI 상호작용 전체에 대한 메타인지적 조절로 확장할 필요가 있다.

이에 본 연구는 AI-SPM을 생성형 AI와의 상호작용 속에서 학습자가 자신의 사고, AI 산출물, 그리고 인간-AI 상호작용 과정을 맥락과 관점의 측면에서 점검하고 조절하는 상위 메타인지 역량으로 개념화하였다. AI-SPM은 메타인지 이론의 지식, 감시, 조절 기능을 바탕으로 하되, 그 적용 대상을 개인 내부의 사고 과정에서 인간-AI 상호작용 전체로 확장한다. 또한 상황인지와 분산인지, 관점주의, 인식론적 CT를 이론적 기반으로 삼아, 상호작용 전, 상호작용 중, 상호작용 후, 장기 축적의 국면에서 작동하는 구성 요소를 체계화하였다. 이를 통해 AI-SPM은 단순한 AI 활용 능력이나 프롬프트 작성 기술이 아니라, AI와 함께 형성되는 사고 과정을 책임 있게 점검하고 조절하기 위한 메타인지적 역량 체계로 제안되었다.

기존 역량 논의와 비교할 때, AI-SPM의 차별점은 점검

과 조절의 대상과 단위에 있다. 자기조절학습이 학습자 내부의 인지 과정을 조절 대상으로, AI 리터러시가 AI 시스템 일반의 신뢰성과 편향을 점검 대상으로 한정하는 것에 비해, AI-SPM은 개별 AI 산출물의 구조와 절차, 산출물과 질문에 내재된 관점, 그리고 인간-AI 상호작용 전반에서의 인지적 위탁 수준을 통합적 조절 대상으로 설정하는 것이다.

본 연구의 교육적 의의는 다음과 같다. 첫째, 정보교육의 목표가 AI를 사용할 수 있는 능력에 머무르지 않고, AI와 함께 사고하는 과정을 점검하고 조절하는 능력으로 확장되어야 함을 제시하였다. 둘째, 생성형 AI 활용 수업은 최종 산출물 중심이 아니라 질문 구성, AI 산출물 검토, 프롬프트 재구성, 판단 유지, 새로운 관점 생성이 드러나는 과정 중심으로 설계될 필요가 있음을 밝혔다. 셋째, 평가는 산출물의 완성도뿐 아니라 학습자가 AI 산출물의 전제와 구조를 어떻게 검토했는지, 어떤 판단을 AI에 위탁하고 어떤 판단을 직접 유지했는지, 최종 판단의 근거를 어떻게 설명하는지를 함께 고려해야 함을 제안하였다. 넷째, AI-SPM은 인식론적 CT를 인간-AI 상호작용 전체의 메타인지적 조절로 확장하는 이론적 틀이라는 점에서 정보교육의 핵심 역량 논의를 새롭게 조직할 수 있는 가능성을 제공한다.

다만 본 연구는 문헌 분석과 이론적 논의를 바탕으로 AI-SPM의 개념과 구성 요소를 제안했다는 점에서 한계를 지닌다. AI-SPM이 실제 수업 맥락에서 어떻게 발현되는지, 학습자의 사고 과정과 학습 결과에 어떠한 영향을 미치는지에 대해서는 후속 실증 연구가 필요하다. 향후 연구에서는 AI-SPM의 구성 요소를 측정할 수 있는 도구를 개발하고, AI-SPM 기반 교수-학습 모형과 평가 루브릭을 구체화할 필요가 있다. 또한 학습자의 발달 수준, 교과 특성, 과제 유형, 문화적 맥락에 따라 AI-SPM이 어떻게 다르게 작동하는지 비교, 분석하는 연구도 요구된다. 이러한 후속 연구를 통해 AI-SPM은 생성형 AI 시대 정보교육에서 인간의 사고 역량을 유지하고 확장하기 위한 실천적 교육 틀로 더욱 정교화될 수 있을 것이다.

참고문헌

- [1] Wing, J. (2006). *Computational thinking*. *Communications of the ACM*, 49(3), 33–35. <https://doi.org/10.1145/1118178.1118215>
- [2] Grover, S., & Pea, R. (2013). Computational thinking in K-12: A review of the state of the field. *Educational Researcher*, 42(1), 38–43. <https://doi.org/10.3102/0013189X12463051>
- [3] Shute, V., Sun, C., & Asbell-Clarke, J. (2017). Demystifying computational thinking. *Educational Research Review*, 22, 142–158. <https://doi.org/10.1016/j.edurev.2017.09.003>
- [4] Weintrop, D., Beheshti, E., Horn, M., Orton, K., Jona, K., Trouille, L., & Wilensky, U. (2016). Defining computational thinking for mathematics and science classrooms. *Journal of Science Education and Technology*, 25(1), 127–147. <https://doi.org/10.1007/s10956-015-9581-5>
- [5] Partnership for 21st Century Skills. (2009). *P21 framework definitions*. Partnership for 21st Century Skills. <https://files.eric.ed.gov/fulltext/ED519462.pdf>
- [6] Allen, L., & Kendeou, P. (2024). ED-AI Lit: An interdisciplinary framework for AI literacy in education. *Policy Insights from the Behavioral and Brain Sciences*, 11(1), 3–10. <https://doi.org/10.1177/23727322231220339>
- [7] UNESCO. (2024). *AI competency framework for students*. United Nations Educational, Scientific and Cultural Organization. <https://www.unesco.org/en/articles/ai-competency-framework-students>
- [8] Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günnemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., ... Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, Article 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- [9] Denny, P., Prather, J., Becker, B. A., Finnie-Ansley, J., Hellas, A., Leinonen, J., Luxton-Reilly, A., Reeves, B. N., Santos, E. A., & Sarsa, S. (2024). Computing education in the era of generative AI. *Communications of the ACM*, 67(2), 56–67. <https://doi.org/10.1145/3624720>
- [10] Brown, J. S., Collins, A., & Duguid, P. (1989). Situated cognition and the culture of learning. *Educational Researcher*, 18(1), 32–42. <https://doi.org/10.3102/0013189X018001032>
- [11] Lave, J., & Wenger, E. (1991). *Situated learning: Legitimate peripheral participation*. Cambridge University Press.
- [12] Hutchins, E. (1995). *Cognition in the wild*. MIT Press.
- [13] Hollan, J., Hutchins, E., & Kirsh, D. (2000). Distributed cognition: Toward a new foundation for human-computer interaction research. *ACM Transactions on Computer-Human Interaction*, 7(2), 174–196. <https://doi.org/10.1145/353485.353487>
- [14] Nietzsche, F. (1974). *The gay science: With a prelude in rhymes and an appendix of songs* (W. Kaufmann, Trans.). Vintage Books. (Original work published 1882)
- [15] Haraway, D. (1988). Situated knowledges: The science question in feminism and the privilege of partial perspective. *Feminist Studies*, 14(3), 575–599. <https://doi.org/10.2307/3178066>
- [16] Hales, S., & Welshon, R. (2000). *Nietzsche's perspectivism*. University of Illinois Press.
- [17] Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, Virtual Event, Canada, 610–623. <https://doi.org/10.1145/3442188.3445922>
- [18] Brennan, K., & Resnick, M. (2012). New frameworks for studying and assessing the development of computational thinking. *Proceedings of the 2012 Annual Meeting of the American Educational Research*

Association, Vancouver, BC, Canada. <https://scratched.gse.harvard.edu/ct/files/AERA2012.pdf>

- [19] Zimmerman, B. J. (2000). Attaining self-regulation: A social cognitive perspective. In M. Boekaerts, P. R. Pintrich, & M. Zeidner (Eds.), *Handbook of self-regulation* (pp. 13–39). Academic Press.
- [20] Flavell, J. H. (1979). Metacognition and cognitive monitoring: A new area of cognitive-developmental inquiry. *American Psychologist*, 34(10), 906–911. <https://doi.org/10.1037/0003-066X.34.10.906>
- [21] Brown, A. L. (1987). Metacognition, executive control, self-regulation, and other more mysterious mechanisms. In F. E. Weinert & R. H. Kluwe (Eds.), *Metacognition, motivation, and understanding* (pp. 65–116). Lawrence Erlbaum Associates.
- [22] Nelson, T., & Narens, L. (1990). Metamemory: A theoretical framework and new findings. In G. H. Bower (Ed.), *The psychology of learning and motivation* (Vol. 26, pp. 125–173). Academic Press. [https://doi.org/10.1016/S0079-7421\(08\)60053-5](https://doi.org/10.1016/S0079-7421(08)60053-5)
- [23] Georgiades, P. (2004). From the general to the situated: Three decades of metacognition. *International Journal of Science Education*, 26(3), 365–383. <https://doi.org/10.1080/0950069032000119401>
- [24] Gerlich, M. (2025). AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), 6. <https://doi.org/10.3390/soc15010006>
- [25] Fan, Y., Tang, L., Le, H., Shen, K., Tan, S., Zhao, Y., Shen, Y., Li, X., & Gašević, D. (2025). Beware of metacognitive laziness: Effects of generative artificial intelligence on learning motivation, processes, and performance. *British Journal of Educational Technology*, 56(2), 489–530. <https://doi.org/10.1111/bjet.13544>



유인환

· 2000 한국교원대학교 컴퓨터교육과(교육학박사)
· 2000~현재 대구교육대학교 교수

✚ 관심분야 : 컴퓨터교육, AI 교육, 피지컬 AI
✉ bluenull@dnue.ac.kr



강민정

· 2013 한국교원대학교 교육학박사
· 2016 대구교육대학교 교육학석사
· 2025~현재 대구교육대학교 교육전문대학원
AI교육 전공 박사과정

✚ 관심분야 : 컴퓨터교육, AI교육, 피지컬 AI
✉ caroline_kmj@naver.com

부 록

〈부록 1〉 AI-SPM의 개념적 도출 과정과 이론적 배경

① 출발점: 생성형AI 환경에서 학습자 역할의 전환

인지 방향의 전환: AI 산출물을 ‘만들어 내는’수행에서, AI 산출물의 의미와 한계를 ‘판단하는’조정으로 역전



② 검토: 기존 핵심 역량 프레임워크

컴퓨팅사고력(CT) 분해, 패턴 인식, 추상화, 알고리즘 설계	4C 창의, 비판, 소통, 협업	AI 리터러시 AI 이해, 활용, 평가	메타인지, SRL 자기 인지, 점검, 조절
--	----------------------	--------------------------	----------------------------



③ 진단: 생성형AI 환경에서 드러난 한계

상황성의 결여 인간-AI 상호작용 체계 반영 못함	관점성의 결여 AI 산출물, 질문의 전제, 관점 점검 부족	통합적 조절의 부재 상호작용 전반에 대한 메타인지적 조절 부재
--------------------------------	--	--



④ 재해석1: CT의 인식론적 확장

동일한 CT 요소(분해, 패턴 인식, 추상화, 알고리즘 설계)가 생성형AI 환경에서 서로 다른 방향으로 작동

실행적CT 학습자가 직접 수행하는 문제 해결 사고 (분해, 패턴 인식, 추상화, 알고리즘 설계)	인식론적CT AI 산출물의 구조, 근거, 추상화 수준, 절차적 타당성을 CT로 판단
---	--



⑤ 재해석2: 메타인지의 확장과 AI-SPM의 고유 기능

메타인지의 점검 대상을 자기 사고에서 산출물, 질문, 맥락, 관점, 위탁 수준으로 확장

사고 출처의 구분 학습자 이해에 근거한 판단과 AI 제안에 의해 형성된 판단의 구분	인지 위탁의 동적 조절 AI에 위탁할 과정과 학습자가 유지할 판단의 실시간 구분	관점, 출처의 감시 산출물, 질문의 전제와 관점, 정보 출처에 대한 감시
--	--	--



⑥ 통합: 이론적 기반의 종합

메타인지 이론 지식, 감시, 조절 기능측	상황인지, 분산인지 인간-AI 상호작용을 전, 중, 후, 장기 축적으로 접근	관점주의 전제, 관점, 능력 점검 근거	인식론적 CT 산출물 구조, 절차 판단
---------------------------	--	--------------------------	--------------------------



AI 맥락 다관점 메타인지(AI-SPM: AI-Situated Perspectival Metacognition)

학습자가 생성형 AI와의 상호작용 속에서 자신의 사고, AI 산출물, 그리고 인간-AI 상호작용 과정을 맥락과 관점의 측면에서 점검하고 조절하는 상위 메타인지 역량

- ‘맥락(Situated)’은 AI 활용이 과제 목적, 평가 기준, 수업 규범, 도구 환경에 따라 달라진다는 점을 의미하며, 학습 상황에서 AI 활용 방식과 위탁 범위가 적절한지를 점검하고 조절하는 역량을 포함
- ‘다관점(Perspectival)’은 AI 산출물의 내재적 관점 점검, 학습자 질문의 관점 성찰, 대안적 관점의 능동적 생성이라는 세 가지 범주를 포함
- ‘메타인지(Metacognition)’는 인간-AI 상호작용 과정 전체를 대상으로 하는 점검과 조절을 의미