

생성형 AI 활용 국어 수업의 문제 양상과 국어과 기반 AI 리터러시 구성 요소 — 국어 교사 FGI 질적 연구

허선아 동래원예고등학교 교사

- I. 연구의 필요성 및 목적
- II. 이론적 배경
- III. 연구 방법
- IV. 연구 결과
- V. 결론 및 교육적 시사점

I. 연구의 필요성 및 목적

생성형 인공지능의 확산은 학습자가 언어를 이해하고 생산하는 방식뿐 아니라 사고를 조직하고 의미를 구성하는 과정에도 변화를 가져오고 있다. 오늘날 학습자는 글쓰기, 말하기, 정보 탐색과 같은 주요 언어 활동에서 인공지능이 생성한 텍스트를 목적과 담화 맥락에 맞게 선택하고 조정해야 한다. 이에 따라 학습자는 출력물을 그대로 수용하는 사용자가 아니라 그 타당성을 검토하고 재맥락화하며 선택·수정·삭제의 근거를 성찰하는 의미 구성자로 기능할 필요가 있다.

그러나 2022 개정 국어과 교육과정은 생성형 인공지능 활용이 본격화되기 이전에 마련된 문서로서 디지털 매체 환경에서의 정보 이해와 비판적 평가 등에 관한 내용은 제시하고 있으나, 생성형 인공지능과의 상호작용 과정에서 요구되는 사실성 검증, 관점 판단, 생성 결과에 대한 책임과 윤리와 같은 판단 과정은 교육과정 문서에서 아직 명시적으로 구조화되어 있지 않다.

AI 리터러시에 관한 논의는 인공지능 기술의 확산과 함께 그 개념과 교육 방향을 정리하려는 다양한 연구로 축적되어 왔다. 일부 연구에서는 AI 교

육의 개념과 범주, 교육과정 구성, 교사의 역할 등을 중심으로 AI 교육의 구조와 과제를 정리하였으며(성은모·김동호·신서경 외, 2023), AI 리터러시의 하위 차원과 구성 요소를 체계화하려는 시도도 이루어졌다(황현정·황용석·박지수 외, 2024). 또한 AI 서비스 이용 경험이나 개인의 혁신성과 같은 요인이 AI 리터러시 형성에 영향을 미친다는 점을 실증적으로 확인한 연구도 있다(정서현·박주연, 2024).

한편 국어교육 및 리터러시 연구에서는 디지털 환경에서의 리터러시를 단순한 도구 활용 능력을 넘어 언어적·사회문화적 의미 구성 능력으로 확장하여 이해하려는 논의가 축적되어 왔다(정현선, 2004). 이러한 관점은 최근 AI 리터러시 논의에서도 기술 이해 중심 접근을 넘어 비판적 사고와 책임, 윤리 등을 포함해야 한다는 주장으로 이어지고 있으며(이유미·박윤수, 2021; 이지영, 2024), 국어 교과 맥락에서 AI 리터러시 교육의 가능성을 탐색하려는 연구도 제시되고 있다(최창원, 2023).

다만 이러한 논의는 생성형 인공지능이 만들어내는 언어와 텍스트를 비판적으로 분석하고 그 맥락과 전제를 성찰해야 한다는 점을 제기하면서도(김운용, 2025: 390-391), 이러한 판단이 교실의 언어 수행 과정에서 어떤 국면에서 이루어지고 어떤 방식으로 약화되는지를 구체적으로 설명하는 데에는 한계가 있다. 특히 생성형 인공지능 환경에서 요구되는 리터러시가 비판적 사고와 책임 있는 사용을 포함해야 한다는 점은 제시되어 왔지만, 이러한 문제가 실제 읽기·쓰기·말하기 수행에서 어떻게 나타나는지에 대한 분석은 충분하지 않다.

실제 수업에서는 학생들이 인공지능 기능을 비교적 능숙하게 활용하면서도 생성된 텍스트를 과제 목적에 맞게 재구성하거나 타당성을 검증하고 선택 근거를 설명하는 과정에서는 어려움을 보이는 장면이 관찰되기도 한다. 이 경우 인공지능 텍스트가 비판적 검토의 대상이 아니라 정답처럼 수용되는 양상이 나타나기도 하며, 이러한 문제는 개별 기능의 부족이라기보다 언어 수행 과정에서 판단이 충분히 작동하지 않는 수행 구조와 관련될 가능

성을 보여준다.

이에 본 연구는 생성형 인공지능 활용 과정에서 나타나는 국어과적 문제를 ‘판단 결손’의 관점에서 살펴보고, 국어 교실 맥락에서 요구되는 AI 리터러시의 구성 요소와 과정 구조를 탐색하고자 한다. 이를 위해 국어 교사 7인을 대상으로 한 포커스 그룹 인터뷰(FGI)와 문헌 분석을 결합하여 교실에서 관찰되는 언어 수행 문제 양상을 분석하고, 이를 토대로 국어과 기반 AI 리터러시의 구성 요소를 도출하고자 한다.

이 연구의 연구 문제는 다음과 같다.

첫째, 생성형 인공지능 활용 과정에서 국어 교사가 인식한 주요 언어 수행 문제 양상은 무엇인가.

둘째, 이러한 문제를 설명하기 위해 국어과 기반 AI 리터러시는 어떤 구성 요소로 체계화될 수 있는가.

셋째, 도출된 구성 요소들은 실제 의미 구성 과정에서 어떠한 방식으로 조직되고 상호 연계되는가.

II. 이론적 배경

1. 생성형 AI 환경과 국어과 언어 수행의 변화

생성형 인공지능의 확산은 국어과에서 이루어지는 언어 수행의 조건을 변화시키고 있다. 기존의 디지털 도구가 자료 탐색이나 정리와 같은 보조적 기능에 주로 활용되었다면, 생성형 AI는 문장과 담화를 직접 산출하며 읽기·쓰기·말하기 수행 과정에 보다 직접적으로 개입한다. 이로 인해 학습자의 언어 수행은 단순한 표현 생성 활동을 넘어 AI가 제시한 텍스트를 검토하고 선택하며 수정하는 판단 과정과 긴밀하게 결합된다.

그러나 이러한 판단 과정이 실제 수업에서 충분히 활성화되는 것은 아닙니다. 생성형 AI 환경에서 학습자는 AI가 산출한 텍스트를 과제 목적과 담화 맥락에 맞게 재구성하고 내용의 타당성을 점검하며 표현 선택의 이유를 설명할 수 있다. 하지만 실제 수업에서는 이러한 검토와 조정이 충분히 이루어지지 않거나 AI 산출물이 별도의 판단 없이 사용되는 경우가 나타난다. 이러한 양상은 생성형 AI 활용이 단순한 도구 사용의 문제가 아니라, 텍스트를 비판적으로 이해하고 판단하며 활용하는 능력과 관련됨을 시사한다. 이유미(2021: 287)는 AI 리터러시를 인공지능 기술을 비판적으로 이해하고 활용하며 변화하는 사회 속에서 주체적으로 살아가기 위한 기초 능력으로 설명하였다.

본 연구에서는 이러한 상황을 ‘판단 결손’이라는 개념으로 설명하고자 한다. 여기서 판단 결손은 학습자의 능력 부족이나 인지적 결함을 의미하는 것이 아니라 수업 과제나 수행 조건 속에서 판단이 개입되어야 할 국면이 충분히 활성화되지 않아 판단 활동이 축소되는 수행 양상을 가리킨다. 예를 들어 학습자가 AI 산출물을 과제 맥락에 맞게 재구성하지 않거나 정보의 타당성을 점검하지 않거나 표현 선택의 이유를 설명하지 않는 경우가 이에 해당한다. 이러한 상황에서는 언어 수행 과정에서 이루어지는 의미 구성 활동이 제한될 수 있다. 실제로 디지털 미디어 리터러시에 대한 학생 인식을 조사한 연구에서도 참여 태도에 대한 인식은 비교적 긍정적으로 나타났으나 문제 해결 과정에서의 실제 참여 태도나 윤리적 태도는 상대적으로 낮게 나타났다(최진영, 2023: 263). 이는 리터러시에 대한 인식이 실제 수행 과정의 판단으로 항상 이어지는 것은 아님을 보여준다.

또한 생성형 AI의 산출 결과에는 오류나 부정확한 정보가 포함될 수 있으며 이를 검토하고 조정하는 인간의 개입이 필요하다(김종규, 2023: 17). 김운용(2025: 340)은 현존 생성형 AI가 감응력과 같은 경험적 기반을 갖지 않기 때문에 도덕적 행위자로 간주하기 어렵다고 지적하였다. 이는 생성형 AI의 언어 산출이 인간과 유사한 형태를 보이더라도 그 자체가 판단이나 책임

의 주체가 될 수 없음을 시사한다.

이러한 논의를 종합하면 생성형 AI 환경에서 국어과 언어 수행은 결과물 생산뿐 아니라 AI가 생성한 텍스트를 검토하고 조정하며 활용하는 판단 과정으로 이해될 필요가 있다.

2. 디지털·AI 리터러시 논의의 한계와 생성형 AI 환경의 특성

디지털 리터러시는 초기에는 컴퓨터와 정보기술을 활용하는 능력을 중심으로 논의되었으나 이후 다양한 정보와 텍스트를 이해하고 평가하며 활용하는 역량을 포함하는 개념으로 확장되어 왔다(김태운·문수연·백지혜 외, 2025: 991). 이는 디지털 환경에서의 문식 활동이 단순한 정보 접근을 넘어 텍스트 해석과 판단을 포함하는 활동으로 이해되고 있음을 보여준다.

최근에는 인공지능 기술의 확산과 함께 AI 리터러시에 대한 논의도 이루어지고 있다. 일부 연구에서는 AI 리터러시를 AI의 작동 원리 이해, 활용 능력, 윤리 인식, 사회적 영향에 대한 이해 등 다양한 요소로 설명하고 있다(최숙영, 2022). 이러한 연구들은 AI 기술이 교육과 사회에 미치는 영향을 이해하는 데 중요한 시사점을 제공한다.

생성형 AI의 등장으로 텍스트 생산 과정에 대한 논의도 새롭게 제기되고 있다. 생성형 AI는 사용자의 요청에 따라 문장과 담화를 생성할 수 있으며 이러한 결과물은 학습자의 읽기와 쓰기 활동에 다양한 방식으로 활용될 수 있다. 이 과정에서 학습자는 AI가 생성한 텍스트를 그대로 사용하는 것이 아니라 내용을 검토하고 수정하거나 재구성하는 활동을 수행할 수 있다. 일부 연구에서는 이러한 과정에서 정보의 신뢰성 판단이나 윤리적 고려가 필요함을 지적하였다(Wang, Rau, & Yuan, 2023; 노환호·김현정·김민진, 2024).

국어교육의 관점에서 이러한 논의는 언어 수행 과정에서 이루어지는 의미 해석과 표현 활동의 문제로 이해될 수 있다. 정현선(2004)은 디지털 리터

러시를 텍스트의 의미를 비판적으로 이해하고 사회문화적 맥락 속에서 해석하는 능력으로 규정하였다. 이러한 관점은 디지털 환경에서 텍스트가 단순한 정보가 아니라 해석과 판단의 대상이 되는 언어 활동임을 강조한다. 한편 AI 리터러시에 관한 논의에서는 이러한 리터러시 개념을 인공지능 환경으로 확장하려는 시도가 이루어지고 있다. 예컨대 이유미(2021: 290)는 AI 리터러시를 단순한 기술 활용 능력이 아니라 AI로 인해 변화하는 문화와 관계를 이해하고 이에 참여할 수 있는 능력으로 설명하였다. 이러한 논의는 AI 환경에서도 텍스트가 단순한 정보가 아니라 해석과 판단이 요구되는 언어 수행의 대상으로 이해될 필요가 있음을 시사한다.

다만 기존 디지털·AI 리터러시 연구는 기술 이해나 활용 능력, 윤리 인식과 같은 범주에 초점을 두는 경향이 있다. 이러한 접근은 생성형 AI가 개입한 언어 수행 상황에서 학습자가 수행하는 판단 과정을 교과 수행의 관점에서 구체적으로 설명하는 데에는 한계를 지닌다. 예컨대 AI가 생성한 텍스트를 과제 맥락에 맞게 수정하거나 정보의 타당성을 점검하거나 표현 선택의 이유를 설명하는 활동은 국어과 수업에서 중요한 수행이지만 이러한 판단 과정은 리터러시 요소로 명확히 제시되지 않았다.

3. 국어과 기반 AI 리터러시의 재구조화 필요성

기존 디지털 리터러시 및 AI 리터러시 연구는 인공지능 기술의 확산에 따라 새로운 매체 환경에서 요구되는 역량을 규정하고 그 구성 요소를 체계화하려는 방향으로 전개되어 왔다. 이러한 논의는 인공지능 기술의 이해, 활용 능력, 윤리적 인식 등 다양한 차원을 포괄하며 AI 환경에서 요구되는 기본 역량을 정리하는 데 중요한 시사점을 제공한다.

그러나 생성형 AI가 직접 텍스트를 산출하는 환경에서는 학습자가 수행하는 언어 활동의 양상이 이전과는 다른 방식으로 나타난다. 생성형 AI는 문장이나 담화를 완결된 형태로 제시하는 경우가 많기 때문에 학습자는 이를

단순히 활용하는 수준을 넘어 텍스트의 적절성을 검토하고 과제 맥락에 맞게 수정하며 그 선택의 이유를 설명하는 과정에 참여하게 된다. 실제 수업에서는 학습자가 AI가 생성한 텍스트를 과제 목적에 맞게 충분히 조정하지 않거나 정보의 사실성과 출처를 점검하지 않거나 표현 선택의 이유를 설명하지 않는 상황이 나타나기도 한다. 이러한 양상은 생성형 AI 환경에서의 언어 수행이 단순한 기술 활용의 문제가 아니라 텍스트를 해석하고 조정하는 판단 과정과 긴밀하게 관련되어 있음을 보여준다.

디지털 리터러시 연구에서는 디지털 환경에서 요구되는 역량을 범주화하고 체계화하려는 논의가 이루어져 왔다. 최재호·전신영(2023: 101-116)은 디지털 리터러시를 지식·기술, 활동, 사고, 시민의식, 태도 등의 범주로 체계화하였다. 이러한 정리는 디지털 환경에서 요구되는 리터러시 역량을 종합적으로 제시하였다는 점에서 의미가 있다.

한편 생성형 AI 환경에 특화된 리터러시 연구도 이루어지고 있다. 김인주·이소율·김귀훈(2024)은 생성형 AI 리터러시를 이해, 활용, 가치의 요인으로 구조화하였으며, 성은모 외(2023)는 AI 사고력과 윤리 교육의 균형을 강조하였다. 또한 정서현·박주연(2024: 160-161)은 AI 서비스 이용 경험과 개인의 혁신성이 AI 리터러시 형성에 영향을 미칠 수 있음을 실증적으로 확인하였다. 이러한 연구들은 생성형 AI 활용과 관련된 리터러시의 구성 요소와 형성 요인을 제시하였다는 점에서 의미가 있다.

그러나 국어과 수업의 관점에서는 이러한 요소들을 단순히 나열하기보다 실제 언어 수행 과정에서 이루어지는 판단을 중심으로 재구성할 필요가 있다. 국어과 수업에서는 학습자가 AI 산출물을 과제 목적과 담화 맥락에 맞게 재구성하고, 정보의 사실성과 관점을 점검하며, 선택한 표현에 대한 설명과 책임을 수행하고, 질문과 응답을 조정하면서 의미를 발전시키는 활동이 중요하게 작동한다. 따라서 생성형 AI 환경에서 요구되는 국어과 기반 AI 리터러시는 언어 수행 과정에서 이루어지는 판단 절차를 중심으로 재구조화될 필요가 있다.

4. 국어과 기반 AI 리터러시의 예비 분석틀

앞 절에서 살펴본 바와 같이 디지털 리터러시와 AI 리터러시 연구는 기술 이해, 활용 능력, 윤리 인식 등 다양한 요소를 중심으로 리터러시 역량을 설명해 왔다. 이러한 논의는 인공지능 환경에서 요구되는 기본 역량을 정리하는 데 중요한 시사점을 제공한다. 그러나 생성형 AI가 직접 텍스트를 산출하는 환경에서는 학습자가 수행하는 언어 활동이 단순한 기술 활용을 넘어 AI 산출물을 검토하고 수정하며 활용하는 과정과 밀접하게 관련된다.

생성형 AI가 제시하는 텍스트는 문장이나 담화 형태의 결과물로 제시되는 경우가 많기 때문에 학습자는 이를 그대로 수용하거나 일부 수정하는 방식으로 과제를 수행하기 쉽다. 이 과정에서 텍스트를 과제 맥락에 맞게 재구성하거나 정보의 사실성과 출처를 점검하거나 표현 선택의 이유를 설명하는 판단 절차가 충분히 이루어지지 않는 장면이 수업에서 나타나기도 한다.

이러한 점을 고려할 때 생성형 AI 환경에서의 국어과 AI 리터러시는 단순한 기술 이해나 활용 능력의 문제로만 설명되기보다는 언어 수행 과정에서 작동하는 판단 활동의 측면에서 재검토될 필요가 있다. 이에 본 연구에서는 선행연구에서 제시된 리터러시 요소들을 국어과 언어 수행의 관점에서 재검토하고 생성형 AI 활용 과정에서 나타나는 판단 국면을 중심으로 예비 범주를 정리하였다.

본 연구에서 주목한 판단 국면은 다음과 같다. 첫째, AI 산출물을 과제 맥락 속에서 해석하고 수정 대상으로 인식하는 판단이다. 둘째, 산출물의 사실성·출처·관점을 점검하는 판단이다. 셋째, 선택하거나 수정한 표현에 대해 발화 주체로서 책임을 인식하고 설명하는 판단이다. 넷째, 질문 수정과 응답 비교를 통해 의미를 조정하는 상호작용적 판단이다.

이러한 판단 국면을 기준으로 선행연구에서 제시된 리터러시 요소를 재정리한 결과 인지·맥락, 비판·검증, 윤리·책임, 상호작용·협상의 네 범주로 정리할 수 있었다. 인지·맥락 범주는 AI 텍스트를 과제 목적과 담화 맥락에

따라 해석하고 재구성하는 인식을 의미하며, 비판·검증 범주는 산출물의 사실성, 출처, 논리, 관점을 점검하는 판단을 포함한다. 윤리·책임 범주는 AI 산출물을 선택하거나 수정하여 사용할 때 그 결과에 대한 발화 책임을 인식하고 설명하는 판단을 가리킨다. 상호작용·협상 범주는 질문 수정과 재질의, 응답 비교를 통해 의미를 조정하는 상호작용 과정을 포함한다.

이러한 범주는 생성형 AI 활용 수업에서 나타나는 수행 양상을 탐색하기 위한 예비 분석틀로 설정된 것이다. 해당 범주는 <표 1>에 제시하였다.

<표 1> 선행연구 기반 AI 리터러시 예비 범주

예비 범주	핵심 의미	국어과 판단 초점
인지·맥락	AI 산출물의 의미와 맥락을 판단하고 과제 상황에 맞게 해석	텍스트를 맥락에 따라 재구성하는 언어적 인식
비판·검증	사실성·출처·관점에 대한 판단	정보의 신뢰성과 논리 구조 검토
윤리·책임	결과물에 대한 언어적 책임 인식	표현 선택과 사용에 대한 발화 책임
상호작용·협상	질문 수정·응답 비교를 통한 의미 조정	AI와의 상호작용 속 의미 협상

<표 1>은 선행연구에서 제시된 디지털·AI 리터러시 요소를 국어과 언어 수행의 판단 과정에 맞추어 재구성한 예비 분석틀이다. 다만 이러한 범주가 실제 수업에서 어떠한 양상으로 나타나는지는 문헌 검토만으로 확인하기 어렵다. 이에 본 연구에서는 국어 교사의 수업 경험을 바탕으로 각 범주가 실제 교실 장면에서 어떻게 인식되는지를 탐색하기 위해 포커스 그룹 인터뷰(FGI)를 설계하였다.

III. 연구 방법

1. 연구 설계 및 절차

이 연구는 II장에서 제시한 국어과 기반 AI 리터러시의 예비 범주를 토대로, 생성형 인공지능 활용 수업에서 교사들이 인식하는 학생 수행의 문제양상을 탐색하기 위해 포커스 그룹 인터뷰(Focus Group Interview, FGI)를 실시하였다. 생성형 AI 환경에서는 학습자가 AI 산출물을 어떻게 이해하고 검토하며 수정하고 사용하는지가 중요한 판단 국면으로 작용한다. 이러한 판단 양상은 실제 수업 맥락에서 다양한 형태로 나타나므로, 교사의 수업 경험을 통해 반복적으로 관찰되는 수행 문제를 탐색할 필요가 있다.

FGI는 참여자 간 상호작용을 통해 교실 사례와 인식을 구체적으로 드러낼 수 있다는 점에서 생성형 AI 활용 수업에서 나타나는 수행 문제를 탐색하는 데 적절한 방법으로 판단하였다. 특히 교사들은 수업 설계, 학생 수행 과정, 평가 경험을 함께 서술할 수 있기 때문에 AI 활용 과정에서 어떤 판단이 약화되거나 생략되는지를 실제 수업 장면을 중심으로 확인할 수 있다.

인터뷰 질문은 <표 1>의 예비 범주를 기준으로 구성하였다. 질문은 학습자가 AI 텍스트를 어떻게 인식하는지, 산출물을 어떤 기준으로 검토하는지, 표현 선택의 책임을 어떻게 인식하는지, AI와의 상호작용 과정에서 의미조정이 어떻게 이루어지는지를 탐색하도록 설계하였다. 범주와 질문의 대응관계는 <표 2>에 제시하였다.

〈표 2〉 FGI 질문 체계

범주	목적	탐색 내용	질문
인지·맥락	AI 텍스트를 판단 대상으로 인식하는 교사의 기준 탐색	맥락 적합성, 과제 목적 인식, AI 이해 수준	• 학생들은 AI 텍스트를 어떤 성격의 텍스트로 인식한다고 보십니까? • 과제 맥락에 맞지 않는 AI 응답을 학생들이 어떻게 받아들이는 경우가 많습니까?
비판·검증	AI 생성 텍스트의 위험 요소에 대한 교사 인식 파악	사실 오류, 출처 불명, 논리 결함, 편향	• 학생들이 AI 텍스트를 사용할 때 가장 빈번하게 나타나는 문제는 무엇입니까? • 사실성이나 출처 문제로 수업에서 실제로 어려움을 겪은 사례가 있습니까?
윤리·책임	AI 활용에서 책임 귀속과 윤리 인식의 양상 탐색	표절 인식, 책임 전가, 선택 이유 설명	• 학생들은 AI가 만든 문장에 대해 책임을 어떻게 인식하고 있다고 보십니까? • AI 활용과 관련해 윤리적 판단이 요구된 수업 장면이 있었습니까?
상호작용·협상	AI-학습자-교사 간 상호작용 구조 분석	프롬프트 수정, 재질의, 의미 협상, 교사 개입	• 학생이 AI와 상호작용하며 의미를 조정하는 장면이 나타난 적이 있습니까? • 프롬프트 수정 과정에서 교사의 언어적 개입은 어떻게 이루어져야 한다고 보십니까?

본 연구에서 설정한 범주는 생성형 AI 활용 과정에서 나타나는 수행 문제를 국어과 언어 수행의 판단 초점에 따라 구분한 것이다. 즉 인지·맥락은 텍스트의 맥락적 해석과 재구성, 비판·검증은 사실성과 논리 점검, 윤리·책임은 표현 선택에 대한 발화 책임 인식, 상호작용·협상은 질문 수정과 응답 비교를 통한 의미 조정을 중심으로 설정하였다.

이러한 예비 범주는 교실에서 관찰되는 수행 문제를 탐색하기 위한 초기 분석틀로 활용되며, 이후 자료 분석 과정에서 실제 수업 맥락에서 나타나는 문제 양상을 중심으로 재구조화된다. 본 연구는 교사들의 수업 경험 속에서 반복적으로 인식되는 수행 문제를 분석함으로써 생성형 AI 환경에서 요구되는 국어과 기반 AI 리터러시의 개념 구조를 탐색하고자 하였다.

2. 연구 참여자 및 자료 수집

이 연구는 생성형 AI 활용 수업 경험이 있는 국어교사 7명을 목적적 표집으로 선정하여, 교실에서 반복적으로 관찰되는 생성형 AI 활용의 문제 양상과 이에 대한 교사 인식을 수집하고자 하였다. 참여자는 생성형 AI 활용 수업 경험, 국어과 수업 전문성, 학교급과 경력의 다양성이 균형 있게 반영되도록 하였다.

선정 기준은 다음과 같다. 첫째, 생성형 AI를 활용한 국어과 수업 운영 경험이 최소 1회 이상 있을 것. 둘째, 읽기·쓰기·말하기 중 최소 2개 이상의 국어과 영역에서 AI 활용 수업을 경험하였을 것. 셋째, 중학교와 고등학교 교사가 모두 포함되어 학교급에 따른 수업 맥락의 차이를 확인할 수 있을 것. 넷째, 교직 경력이 특정 구간에 편중되지 않도록 경력 초기·중기·후기 교사를 고르게 포함할 것. 반면 생성형 AI 활용이 단순 시연이나 일회적 체험 수준에 그친 경우, 또는 국어과 수업과 직접적인 연계 없이 행정·연수 차원에서만 AI를 접한 경우는 참여 대상에서 제외하였다. 연구 참여자의 기본 정보는 <표 3>과 같다.

<표 3> 연구 참여자 구성

구분	전공 및 학력	교직 경력	AI·국어과 수업 관련 경험
교사 1	독서교육 전공 박사과정 수료	중·고등학교 8년	AI 기반 쓰기 지도, 문식성 이론 적용 경험
교사 2	화법교육 전공 석사과정 수료	고등학교 31년	토론·자료 탐색 수업에서 AI 활용 다수 경험
교사 3	학사	중학교 15년	디지털·AI 활용 읽기 수업 운영
교사 4	학사	중학교 8년	생성형 AI 활용 수행평가 운영
교사 5	학사	중학교 6년	프롬프트 기반 글쓰기 지도 경험
교사 6	교육공학 전공 석사	고등학교 10년	AI 교과서 및 교육자료 개발 연구 참여
교사 7	화법교육 전공 석사	중학교 11년	생성형 AI 활용 프로젝트 수업 운영

자료 수집은 반구조화 질문지를 활용한 포커스 그룹 인터뷰(FGI) 방식으로 이루어졌다. 인터뷰는 연구자가 진행을 맡아 질문 제시와 토의 조정을 수행하였으며, 참여 교사들이 실제 수업 사례를 중심으로 경험을 공유하도록 유도하였다. 질문지는 <표 2>에 제시된 총 8문항으로 구성되었으며, 생성형 AI 활용 수업 경험이 있는 국어교사 3인의 전문가 검토를 거쳐 문항의 명확성, 중복 여부, 현장 적합성을 중심으로 수정·보완하였다.

질문지는 인터뷰 시작 시 참여 교사에게 인쇄물로 제공되었으며, 토의 흐름에 따라 질문 순서는 유연하게 조정하였다. 인터뷰는 약 90분 동안 진행되었고, 모든 발화는 참여자의 동의를 얻어 녹음한 뒤 전사하였다. 전사 과정에서는 의미 해석에 직접적인 영향을 미치지 않는 웃음, 침묵, 중복 발화는 간략히 정리하였고, 판단이나 인식과 관련된 발화는 가능한 한 원문을 유지하였다.

분석 과정에서 여섯 번째와 일곱 번째 참여자의 인터뷰 시점부터 새로운 범주나 핵심 코드의 출현 빈도가 현저히 감소하고 기존 범주의 하위 요소가 반복적으로 진술되는 양상이 확인되었다. 이러한 범주 포화의 징후를 고려할 때, 본 연구의 탐색 목적에 비추어 참여자 수는 충분한 수준이라고 판단하였다.

3. 자료 분석 방법 및 신뢰성 확보 절차

FGI 전사 자료는 개방 코딩, 축 코딩, 선택 코딩의 절차에 따라 분석하였다. 먼저 개방 코딩 단계에서는 교사 발화를 의미 단위로 분절하여 생성형 AI 활용 과정에서 나타나는 학생의 수행 문제와 판단 양상에 주목하며 1차 코드를 도출하였다. 이때 발화가 나타난 수행 맥락을 고려하여 기술적 표현보다는 의미 중심의 코드로 정리하였다.

축 코딩 단계에서는 도출된 1차 코드를 비교·대조하여 공통된 의미 축을 중심으로 재배열하였다. 이 과정에서 II장에서 제시한 예비 범주와의 관

런성을 함께 검토하였다. 범주화의 기준은 문제 양상이 국어과 언어 수행의 어느 판단 지점에서 드러나는가에 두었다. 예를 들어 과제 목적과 답화 맥락에 맞게 텍스트를 재구성하지 못하는 문제는 인지·맥락 범주로, 사실성·출처 점검이 이루어지지 않는 문제는 비판·검증 범주로 분류하였다. 또한 표현 선택의 책임 인식이 약화된 발화는 윤리·책임 범주로, 질문 수정이나 응답 비교를 통한 의미 조정이 나타나지 않는 양상은 상호작용·협상 범주로 정리하였다. 하나의 발화에서 여러 문제 양상이 나타나는 경우에는 해당 발화가 드러내는 판단 초점을 기준으로 범주를 결정하였다.

분석 과정의 추적 가능성을 확보하기 위해 원자료에서 1차 코드와 범주로 통합되는 코딩 과정을 <표 4>에 예시적으로 제시하였다.

<표 4> 원자료-코드-범주의 코딩 과정 예시

원자료(교사 발화)	1차 코드	상위 범주	핵심 개념
“학생들이 요약이나 재작성은 잘하는데, 그걸 과제 목적에 맞게 다시 조정하는 건 거의 안 해요.”	과제 목적 고려 부족	인지·맥락	맥락적 판단 약화
“그렇듯하면 그냥 쓰고, 이게 맞는지 확인하는 과정은 잘 안 나와요.”	사실 검증 생략	검증 문제	사실성 판단 부족
“왜 이 문장을 썼냐고 물으면 ‘시가 그렇게 잦아요’라고 답하는 경우가 많아요.”	발화 책임 회피	윤리·책임	표현 선택 책임 약화
“프롬프트를 수정하긴 하지만, 의미 조정보다는 더 나은 AI 답변을 반복적으로 탐색하는 경우가 많아요.”	상호작용 조정 부족	상호작용·협상	의미 협상 약화

<표 4>는 교사 발화를 바탕으로 1차 코드와 범주가 형성되는 과정을 예시적으로 제시한 것이다. 서로 다른 문제 발화들은 축 코딩과 선택 코딩을 거치면서 AI 생성 텍스트를 판단의 대상으로 충분히 검토하지 않는 수행 양상으로 수렴하였다.

선택 코딩 단계에서는 상위 범주 간의 관계를 검토하여 국어과 기반 AI 리터러시의 개념 구조를 정리하였다. 분석 결과, 인지·맥락, 비판·검증, 윤

리·책임, 상호작용·협상 범주는 생성형 AI 활용 과정에서 연쇄적으로 작동하는 판단 국면으로 이해되었다.

분석의 신뢰성을 확보하기 위해 연구자 메모를 활용하여 코딩 기준과 해석 근거를 지속적으로 점검하였다. 또한 국어교육 전공 박사과정 수료 연구자 2인과 교차 검토를 실시하여 범주 명명과 해석의 적절성을 검토하였다. 교차 검토 결과 약 85% 이상의 해석 일치를 확인하였으며, 일부 코드의 분류 기준을 보완하였다. 아울러 참여 교사에게 일부 전사 자료와 초기 범주 해석을 공유하여 참여자 확인을 실시하였다.

IV. 연구 결과

1. 교사 FGI로 본 생성형 AI 활용의 문제 양상

국어과 수업에서 생성형 AI 활용이 빠르게 확산되는 가운데, 교사들은 학생들의 어려움을 단순한 기술 조작 능력의 부족이라기보다 AI 텍스트를 다루는 과정에서 요구되는 언어적 판단이 충분히 작동하지 않는 문제로 인식하고 있었다. 이는 생성형 AI 텍스트가 불확실성과 편향 가능성을 내포한다는 선행 논의(이지영, 2024)와 문제의식을 공유하면서도, 이러한 판단의 약화가 읽기·쓰기·말하기 수행 장면에서 구체적으로 드러난다는 점에서 의미가 있다.

FGI 분석 결과, 교사 진술은 생성형 AI 활용 과정에서 요구되는 판단 지점에 따라 네 가지 문제 양상으로 정리되었다. 이는 학습자가 AI 텍스트를 어떻게 이해하고, 내용을 어떻게 점검하며, 선택한 표현의 책임을 어떻게 인식하고, 상호작용 속에서 의미를 어떻게 조정하는가와 관련된다. 실제 수업에서는 이러한 문제 양상이 동시에 나타나기도 했으나, 본 연구에서는 교사 발

화에서 강조된 판단 지점을 기준으로 범주를 구분하였다.

1) 인지·맥락 판단 과정의 악화

인지·맥락 판단 과정의 악화는 AI 산출물의 사실 여부를 점검하지 않는 상황이라기보다, 그 이전 단계에서 해당 텍스트를 과제 목적, 장르, 청자, 경험 맥락에 맞게 다시 구성해야 할 대상으로 인식하지 못하는 상황을 가리킨다. FGI에 참여한 교사들은 학생들이 질문 입력, 형식 변환, 분량 조정 등 생성형 AI의 도구적 기능 자체는 비교적 능숙하게 활용하지만, AI 산출물을 판단과 조정의 대상이 되는 텍스트로 인식하지 않는 양상이 반복적으로 나타난다고 설명하였다. 교사들의 진술에 따르면 문제의 핵심은 기술 숙련의 부족이 아니라, 산출된 텍스트의 의미·근거·표현이 과제 목적과 맥락에 적합한지 검토하는 인지·맥락 판단 과정이 충분히 작동하지 않는 데 있었다. 이러한 인식은 참여 교사들의 다양한 수업 경험에서 공통적으로 확인되었다.

교사들은 학생들이 AI 산출물을 ‘초안’이나 ‘언어 자원’으로 다루기보다, 이미 완성된 답안처럼 받아들이는 경향이 있다고 설명하였다. 한 교사는 이를 다음과 같이 언급하였다.

“요즘 학생들은 기술이 부족한 게 아니에요. 오히려 제가 모르는 기능을 더 잘 찾아서 쓰기도 해요. 그런데 AI가 써 준 문장을 보면, ‘이게 왜 맞는지, 왜 틀렸는지, 이 표현이 왜 어색한지’를 생각 자체를 안 해요. 그냥 결과를 복사해서 붙이는 게 끝이에요.”
(교사 1)

다른 교사들 역시 “AI가 틀릴 수 있다는 감각 자체가 약하다”는 점을 강조하며, 학생들이 AI 출력을 교과서나 모범답안과 유사한 권위를 지닌 텍스트로 인식하는 경향을 언급하였다(교사 3, 교사 4). “왜 이 표현을 선택했는지, 이 내용이 사실인지”를 묻는 질문에 “그냥 AI가 이렇게 써 줘서요”라는 응답이 반복된다는 진술은, 학생들이 표현 선택과 제출 행위를 자신의 판단

결과라기보다 AI가 제공한 결과를 수용하는 방식으로 이해하는 경향을 보여 준다.

또한 교사들은 학생들이 AI 문장을 과제 목적이나 개인 경험 맥락에 맞게 재구성하지 못하는 사례를 여러 수업에서 관찰하였다고 설명하였다. 예를 들어 진로 소개 글쓰기 수업에서는 다음과 같은 장면이 나타났다고 한다.

“AI로 글을 생성해서 가져온 학생이 있었는데, 겉으로 보면 문장도 매끄럽고 내용도 그럴듯해요. 그런데 읽어 보니까 그 학생이 평소 말하던 진로랑 완전히 달라요. ... ‘그냥 AI가 이렇게 써 줘서요’라고만 해요. 자기 글이라고 인식하지 않으니까, 어디를 바꾸고 왜 바뀌어야 하는지도 모르겠다고 하더라고요.” (교사 3)

토론 준비 과정에서도 학생들은 자신의 주장과 직접 연결되지 않는 근거라 하더라도 “AI가 찾아주었다”는 이유로 그대로 사용하는 경향을 보였으며, 근거와 주장 간의 관계를 다시 조정하기보다 제시된 문장을 유지하려는 태도가 나타났다고 교사들은 설명하였다(교사 2). 더 나아가 감상문, 설명문, 주장문 등 과제 장르와 무관하게 유사한 형식의 문장을 반복적으로 차용하면서 장르 규범이나 수행 목적이 충분히 고려되지 않는 사례도 언급되었다(교사 5).

다만 일부 교사들은 초안의 즉시 제출을 제한하고 수정이나 삭제의 이유를 언어화하도록 요구한 과제 조건에서 학생들이 AI 산출물을 판단 대상으로 인식하는 장면이 나타났다고 설명하였다(교사 6). 이러한 사례는 텍스트를 재맥락화하는 판단 과정이 학습자의 고정된 능력이라기보다 교사의 개입이나 과제 설계, 평가 구조에 따라 활성화될 수 있음을 보여 준다.

종합하면 생성형 AI 활용 수업에서 나타나는 문제는 기술 활용 능력의 부족이라기보다 AI 산출물을 과제 수행 맥락 속에서 다시 해석하고 조정해야 할 텍스트로 인식하는 판단 과정이 충분히 작동하지 않는 상황과 관련된다. 이는 인지·맥락 판단이 적절한 수업 설계와 평가 조건 속에서 형성되고

가시화될 수 있는 교육적 과제임을 시사한다.

2) 정보 타당성 검토 판단의 축소

정보 타당성 검토 판단의 축소는 AI로 산출한 내용의 사실성, 출처, 관점, 논리적 타당성을 점검하는 판단이 충분히 이루어지지 않는 양상에 초점을 둔다. FGI에 참여한 교사들은 생성형 AI 활용 수업에서 학생들이 AI 산출물을 검증할 전제로 읽어야 할 텍스트로 인식하지 않는 경향을 공통적으로 설명하였다. 학생들은 문장의 유창성이나 설명의 완결성을 신뢰의 근거로 삼아 내용을 수용하는 경우가 많았으며, 출처 확인이나 사실성 점검으로 이어지는 정보 검토 판단 과정은 충분히 활성화되지 않는 경우가 많았다. 그 결과 AI 텍스트는 비판적 독해의 대상이라기보다 과제 수행을 빠르게 완결하기 위한 설명으로 활용되는 경향을 보였다. 이러한 인식은 참여 교사들의 다양한 수업 사례에서 공통적으로 확인되었다.

이러한 문제는 조사·설명 중심 과제에서 특히 두드러졌다. 자료 조사 보고서나 개념 설명 글쓰기 수업에서 학생들은 AI가 제시한 정보를 그대로 옮겨 적으면서도, 해당 정보의 출처나 확인 가능성을 질문받을 때 “AI가 정리해 준 내용”이라는 이유로 설명을 중단하는 경우가 많았다. 교사 5는 조사 보고서 점검 장면을 다음과 같이 설명하였다.

“자료 출처를 쓰라고 했는데, AI로 정리한 문단에는 출처가 하나도 없었어요. ‘이 정보 어디서 온 거니?’라고 물으니까 ‘AI가 써 준 거예요’라고만 해요. 자연스럽게 설명돼 있으니까 사실이라고 믿는 거죠.” (교사 5)

수행평가 상황에서도 유사한 양상이 나타났다. 발표 자료 점검 과정에서 AI 산출물의 오류가 그대로 포함된 사례가 있었으며, 교사가 오류를 지적하자 학생은 정보의 진위를 스스로 확인하기보다 “AI가 알려준 내용”이라는 이유로 추가적인 검토를 시도하지 않는 반응을 보였다. 교사 4는 다음과 같

이 설명하였다.

“발표 슬라이드를 보니까 연도랑 용어가 틀렸어요. 그래서 ‘이건 확인이 필요하다’고 했더니 학생이 ‘AI가 알려준 거예요’라고 말해요. 왜 틀렸는지, 어디서 확인해야 하는지보다 ‘AI가 썼다’는 말로 판단을 끝내요.” (교사 4)

이 장면은 학생들이 오류를 발견하고 교차 확인해야 할 필요를 충분히 인식하지 못한 채, 검증 절차를 더 이상 수행하지 않는 상황을 보여 준다. 교사들의 진술에 따르면 학생들에게 문제는 ‘확인 방법을 모른다’기보다 확인을 수행해야 한다는 절차적 인식이 충분히 활성화되지 않는 데 있었다.

관점과 편향에 대한 판단 역시 사회·문화 주제를 다루는 수업에서 두드러지게 나타났다. 토론 준비나 주장문 쓰기에서 AI가 특정 입장에 유리한 서술을 제공했음에도 학생들은 이를 하나의 관점으로 인식하기보다 곧바로 사실로 받아들이는 경향을 보였다. 교사 6은 토론 자료 준비 장면을 다음과 같이 설명하였다.

“사회 이슈 주제로 토론 준비를 시키면 AI가 한쪽 입장에서 정리해 주는 경우가 많아요. 그런데 학생들은 그걸 ‘한 관점’으로 보지 않아요. ‘다른 시각도 찾아봤니?’라고 물으면 ‘AI가 이렇게 말했어요’로 끝나요.” (교사 6)

이러한 사례에서 학생들은 주장과 근거의 적합성이나 관점의 편향을 점검하기보다 AI 서술의 정돈된 형식을 신뢰의 기준으로 삼는 경향을 보였다. 그 결과 사실과 의견의 구분이 흐려지고, 근거 비교나 대안 검토와 같은 논증 판단은 충분히 이루어지지 않는 경우가 나타났다.

다만 일부 교사들은 출처 대조를 수행평가의 필수 조건으로 제시하거나 AI 응답과 외부 자료를 병렬적으로 비교하도록 설계한 수업에서 학생들이 검증 질문을 스스로 생성하는 장면이 나타났다고 설명하였다(교사 2). 이는

비판·검증 과정 역시 학습자의 고정된 능력이라기보다 과제 구조와 교사 개입에 따라 활성화될 수 있는 수행 절차임을 보여 준다.

종합하면 학생들은 생성형 AI 산출물을 검증의 대상이라기보다 신뢰 가능한 설명으로 전제하는 경향을 보였으며, 그 결과 국어과적 비판·검증 절차가 충분히 작동하지 않는 경우가 나타났다. 이러한 양상은 단순한 정보 활용 기술의 문제가 아니라 읽기 판단 과정이 충분히 활성화되지 않는 상황과 관련된 것으로 이해할 수 있다.

3) 표현 선택과 책임 인식 판단의 약화

표현 선택과 책임 인식 판단의 약화는 AI 산출물의 오류를 검증하지 못하는 상황 자체보다, 선택·수정·제출한 표현을 자신의 언어 행위로 인식하고 그 선택 이유를 설명하지 못하는 양상에 초점을 둔다. 교사들은 학생들이 AI가 생성한 문장을 수행물에 포함시키는 순간에도 그 표현의 의미와 결과를 자신의 언어 행위로 인식하지 않는 경향을 반복적으로 지적하였다. 2)에서 드러난 문제가 AI 산출물을 판단의 대상으로 인식하지 않는 상황과 관련된다면, 이 항에서 확인되는 문제는 판단 이후 단계에서 선택한 표현에 대한 책임과 근거가 충분히 언어화되지 않는 상황과 관련된다. 이러한 인식은 발화 주체성과 관련된 문제로서 참여 교사들의 진술에서 공통적으로 나타났다.

FGI에 참여한 교사들은 학생들이 생성형 AI 활용 과정에서 ‘생성한 주체’와 ‘선택한 주체’를 분리하여 이해하는 경향이 있다고 설명하였다. 학생들은 자신이 문장을 선택하여 제출했다는 사실보다 그 문장을 “누가 만들었는가”에 책임을 귀속시키며 언어 행위의 주체성을 AI 쪽으로 이동시키는 경향을 보였다. 한 교사는 수행평가 피드백 상황을 다음과 같이 설명하였다.

“AI 문장을 그대로 가져와 쓴 학생에게 표현이 부적절하다고 지적하면, 바로 ‘그건 제가 쓴 게 아니라 AI가 쓴 거예요’라고 말해요. 마치 책임이 자동으로 빠져나

가는 느낌이에요. 자기가 선택해서 제출했다는 인식이 거의 없어요.” (교사 2)

이 진술은 학생들이 언어 행위의 책임을 ‘생성’에만 귀속시키고 ‘선택’의 책임을 충분히 인식하지 못하는 상황을 보여 준다. 국어과에서 발화의 책임은 ‘누가 처음 만들었는가’보다 ‘누가 그 말을 채택하고 공적으로 제시했는가’에 귀속되지만, 생성형 AI 활용 상황에서는 이러한 원리가 충분히 작동하지 않는 경우가 나타났다.

이러한 책임 귀속 문제는 말하기 수행에서도 유사하게 나타났다. 발표 수업에서 AI로 작성한 발표문을 활용한 학생이 내용 오류나 표현의 부적절성을 지적받았을 때 스스로 수정하거나 설명하기보다 책임을 AI에 전가하는 반응이 반복적으로 나타났다고 교사들은 설명하였다.

“발표 내용이 애매하다고 지적하면, 학생이 ‘이 부분은 AI가 써 준 거라서요’라고 말해요. 그런데 그 발표는 분명히 그 학생이 한 말이잖아요. 자기 발화에 대한 책임감이 확실히 약해졌다고 느껴요.” (교사 7)

출처와 표절 문제에서도 유사한 양상이 나타났다. 학생들은 AI 문장의 출처 불분명성을 자신의 윤리적 판단 문제로 인식하기보다 AI의 산출 특성으로 설명하며 책임 인식을 충분히 연결하지 못하는 경향을 보였다.

“출처를 물어보면 ‘AI가 만든 거라서 출처가 없어요’라고 해요. 그런데 그걸 쓰기로 결정한 건 학생이잖아요. 그 선택이 책임이라는 걸 잘 연결하지 못해요.” (교사 1)

책임 귀속의 약화는 성찰 단계에서도 이어졌다. 교사들은 학생들이 AI 활용 이후 성찰을 요구받을 때 자신의 판단과 선택을 되돌아보기보다 도구 사용의 편의성에 대한 소감 수준에 머무르는 경우가 많다고 설명하였다.

“성찰지를 쓰게 하면 ‘시간이 절약됐다’, ‘편했다’는 말은 많은데, 왜 이 문장을 선택했는지, 어떤 기준으로 수정했는지는 거의 안 나와요. 선택의 책임을 돌아보는 글이 아니에요.” (교사 1)

이러한 양상은 성찰이 판단과 책임을 언어적으로 재구성하는 과정이라기보다 도구 활용 경험을 평가하는 활동으로 축소되는 경향이 있음을 보여준다. 종합하면 학생들은 AI 산출물을 판단하는 과정뿐 아니라 판단 이후 선택한 언어에 대한 책임을 충분히 자신의 언어 행위로 인식하지 않는 경우가 나타났으며, 이는 국어과 기반 AI 리터러시가 단순한 활용 능력을 넘어 선택한 말에 대한 책임을 인식하고 그 근거를 언어화하는 책임 귀속 및 성찰 역할을 포함할 필요가 있음을 시사한다.

4) 의미 조정 및 상호작용 판단의 약화

의미 조정 및 상호작용 판단의 약화는 결과물의 타당성이나 책임 귀속 여부와는 별개로, 재질문, 조건 수정, 응답 비교를 통해 의미를 조정하는 대화적 판단 과정이 충분히 형성되지 않는 상황을 가리킨다. 교사 FGI 분석에서 반복적으로 언급된 문제는 학생들의 생성형 AI 활용이 질의와 응답의 단회적 구조에 머무르며 의미를 점진적으로 조정하거나 협상하는 상호작용 과정으로 확장되지 않는다는 점이었다. 이러한 인식은 참여 교사들의 진술에서 공통적으로 나타났으며, 생성형 AI와의 상호작용이 과정적 탐색이라기보다 결과를 빠르게 얻기 위한 요청으로 활용되는 경향을 보여 준다.

FGI에 참여한 교사들은 학생들이 생성형 AI를 사고를 확장하는 상호작용의 파트너라기보다 결과를 신속히 제공하는 도구로 인식하는 경향이 있다고 설명하였다. 그 결과 첫 응답 이후 상호작용이 종료되고, 응답의 적합성을 재검토하거나 질문을 재구성하려는 시도는 비교적 드물게 나타났다.

“아이들이 AI를 쓸 때 보면, 질문을 딱 한 번 던지고 답을 받으면 거기서 끝나요.

그 답이 과제랑 조금 어긋나 있어도, ‘왜 이런 답이 나왔을까?’를 다시 묻지 않아요. 제가 ‘그럼 조건을 바꿔서 다시 물어보면 어떨까?’라고 말해야 그제야 다시 입력을 해요. 스스로 대화를 이어 간다는 느낌은 거의 없어요.” (교사 4)

이러한 양상은 특히 쓰기 수업에서 두드러졌다. 학생들은 AI 응답이 추상적이거나 논지와 어긋난다는 점을 인식하면서도 이를 해결하기 위한 질문의 재구성에는 어려움을 보였다.

“예를 들어 주장문 초안을 AI로 만들었는데, 내용이 너무 두루뭉술하다고 느끼는 학생이 있어요. 그래서 ‘이 글의 주장이 뭐야?’라고 물으면 학생도 ‘좀 흐려요’라고는 해요. 그런데 그다음에 AI에게 다시 묻는 질문은 처음 질문이랑 똑같아요. ‘좀 더 명확하게 써 줘’ 같은 말은 해도, ‘근거를 세 가지로 나눠서’, ‘반대 입장을 먼저 제시하고 반박해 줘’처럼 구체적인 조정 질문은 거의 안 나와요.” (교사 6)

토론 준비와 같이 비교와 대안 탐색이 요구되는 활동에서도 상호작용의 단계는 반복적으로 드러났다. 교사들은 복수의 응답을 생성하도록 지도했음에도 이후의 비교가 의미 차이나 논증 구조에 대한 협상으로 이어지지 않는 경우가 많다고 설명하였다.

“토론 준비할 때 일부러 ‘AI에게 서로 다른 관점으로 세 번 물어보자’고 시켜요. 그런데 답을 세 개 받아 놓고 비교하라고 하면, ‘이게 더 길어서 좋아요’, ‘말이 더 쉬워요’ 정도예요. 어떤 관점이 강화됐는지, 근거가 얼마나 설득력 있는지로 선택하지는 않아요.” (교사 3)

말하기 수업에서도 유사한 양상이 나타났다. 발표문 수정 과정에서 학생들은 청중·상황·목적을 반영한 조건 협상적 질문을 구성하기보다 단순

재요청에 머무르는 경우가 많았다.

“발표문을 다듬으면서 ‘청중이 중학생이면 어떻게 달라질까?’라고 물어보라고 하면, 아이들이 다시 AI에게 묻긴 해요. 그런데 질문을 보면 그냥 ‘발표문 다시 써 줘’예요. 청중 조건을 명확히 넣거나, 말하기용 표현으로 바꿔 달라는 협상 질문은 거의 없어요.” (교사 7)

교사들은 이러한 양상의 배경을 학생들이 질문을 의미를 조정하는 사고의 장치로 인식하지 못한다는 점에서 찾았다. 재질문은 오류 수정이나 관점 확장을 위한 협상이라기보다 동일한 요구의 반복으로 나타났으며, AI 응답의 한계를 드러내는 비판적 질문 역시 비교적 드물게 나타났다.

다만 일부 교사들은 과제 구조와 평가 조건이 명확하게 제시된 수업에서는 학생들이 AI 응답 이후에도 재질문과 조정을 수행하는 장면이 나타났다고 설명하였다. 예컨대 질문 조건의 변화를 요구한 경우 학생들은 응답을 비교하며 질문을 수정하였으며, 이는 교사 개입과 평가 기준이 상호작용을 단회적 요청이 아니라 조정 과정으로 확장시키는 조건으로 작동했음을 보여준다.

종합하면 상호작용 과정의 제한은 학습자의 고정된 능력 문제라기보다 과제 설계와 피드백 조건에 따라 달라지는 수행 과정으로 이해할 수 있다. 이러한 결과는 국어과 기반 AI 리터러시가 결과 산출 능력뿐 아니라 의미를 점진적으로 조정하는 대화의 과정성을 핵심 요소로 포함할 필요가 있음을 시사한다.

이상에서 살펴본 네 가지 문제 양상은 서로 독립적으로 발생한다기보다 생성형 AI를 활용한 언어 수행 과정의 서로 다른 판단 지점에서 나타나는 수행상의 결손으로 이해할 수 있다. 즉 인지·맥락 판단의 약화는 AI 산출물을 과제 맥락 속에서 재구성해야 할 텍스트로 인식하지 못하는 상황과 관련되며, 비판·검증 과정의 약화는 해당 텍스트의 사실성, 출처, 관점, 논리적 타

당성을 점검하는 판단 절차가 충분히 작동하지 않는 상황과 관련된다. 또한 책임 귀속 판단의 약화는 선택·수정·제출한 표현을 자신의 언어 행위로 인식하지 못하는 문제와 연결되며, 상호작용·협상 구조의 약화는 재질문과 조건 수정, 응답 비교를 통해 의미를 점진적으로 조정하는 대화적 과정이 충분히 형성되지 않는 상황과 관련된다. 이러한 결과는 생성형 AI 활용 상황에서 나타나는 문제를 단순한 기술 활용 능력의 부족으로 설명하기보다, AI 산출물을 이해하고 검토하며 선택하고 조정하는 언어적 판단 과정의 구조와 관련된 교육적 과제로 이해할 필요가 있음을 시사한다.

2. 국어과 기반 AI 리터러시 구성 요소

교사 FGI 분석 결과, 생성형 AI 활용 과정에서 나타나는 학생들의 어려움은 기술 조작 능력의 부족이라기보다 언어 수행의 핵심 국면에서 요구되는 판단이 약화되거나 중단되는 수행 구조로 나타났다. 학생들은 AI를 활용하여 텍스트를 생성하거나 형식을 변환하는 기능적 활용에는 비교적 익숙하였으나, 그 결과를 자신의 언어 행위로 전환하는 과정에서는 판단과 책임이 충분히 작동하지 않는 경향을 보였다. 이는 생성형 AI 환경에서 국어과 언어 활동의 핵심 과정이 자동적으로 유지되지 않음을 시사한다.

본 연구는 이러한 문제 양상을 설명하기 위한 출발점으로 문헌 분석을 통해 인지·맥락, 비판·검증, 윤리·책임, 상호작용·협상의 네 예비 범주를 설정하였다(〈표 1〉). 이후 교사 FGI 분석을 통해 교실 맥락에서 반복적으로 나타나는 수행 문제를 기준으로 범주의 의미와 기능을 재검토하여 국어과 기반 AI 리터러시 구성 요소를 재구조화하였다.

분석 결과 예비 범주는 다음과 같은 수행 중심 개념으로 구체화되었다. 먼저 인지·맥락 범주는 AI 산출물을 과제 목적과 담화 맥락에 맞게 재구성하지 못하는 문제로 나타났으며, 이는 텍스트를 수행 맥락 속에서 다시 조직하는 재맥락화 수행으로 정리되었다. 둘째, 비판·검증 범주는 사실성·출

처·관점 점검이 이루어지지 않는 문제로 나타났으며, 이는 정보의 타당성을 판단하는 검증 수행으로 구체화되었다. 셋째, 윤리·책임 범주는 표현 선택의 책임을 AI에 전가하는 양상으로 나타났으며, 이는 발화 주체성을 명확히 하는 책임 귀속 수행으로 재정의되었다. 마지막으로 상호작용·협상 범주는 단회적 질의-응답 구조에 머무르는 문제로 나타났으며, 이는 질문 수정과 응답 비교를 통해 의미를 조정하는 의미 협상 수행으로 정리되었다.

이러한 관계는 교실에서 확인된 문제 양상과 대응 요소를 중심으로 <표 5>에 제시하였다.

<표 5> 교실에서 확인된 문제 양상과 대응 구성 요소

교실에서 확인된 문제 양상	문제의 성격	대응요소
AI 텍스트를 완성된 답안으로 수용하고, 목적·장르·경험에 따른 재구성에 실패함	인지·맥락 판단 과정의 약화	재맥락화
출처·사실성·관점에 대한 검증 없이 AI 응답을 신뢰함	정보 타당성 검토 판단의 축소	검증
“AI가 했으니 나는 아니다”라는 책임 회피로 발화 주체성이 약화됨	표현 선택 판단의 생략	책임 귀속
단발성 질의-응답에 머물며 의미 조정·협상이 이루어지지 않음	의미 조정 및 상호작용 판단의 약화	의미 협상

<표 5>는 교사들이 인식한 문제들이 개별 행동의 오류라기보다 언어 수행 과정에서 판단이 약화되는 구조적 양상임을 보여준다. FGI 분석 결과 네 문제 양상은 독립적으로 발생하기보다 서로 연쇄적으로 나타나는 경향을 보였다. 예를 들어 재맥락화가 이루어지지 않을 경우 사실성 검증이 뒤따르지 않았으며, 검증이 생략된 상태에서는 표현 선택에 대한 책임 인식도 약화되는 경향이 나타났다.

이에 본 연구는 국어과 기반 AI 리터러시를 생성형 AI 활용 상황에서 언어 수행의 판단과 책임을 회복하는 능력으로 정의하고, 그 구성 요소를 재맥

락화, 검증, 책임 귀속, 의미 협상의 네 범주로 정리하였다. 이하에서는 각 요소를 수행 지표 중심으로 구체화하고 수업 및 평가로의 적용 가능성을 논의한다.

1) 재맥락화

재맥락화는 생성형 AI가 산출한 문장을 완성된 답으로 그대로 사용하는 것이 아니라, 과제 목적과 답화 맥락, 청자, 자신의 경험 등을 고려하여 내용을 다시 검토하고 조정하는 수행 능력을 의미한다. 즉 학습자가 AI가 제시한 문장을 초안 텍스트로 인식하고, 과제의 요구와 글의 흐름에 맞게 수정하거나 재구성하는 과정이 포함된다. 이러한 개념은 예비 범주 가운데 인지·맥락 범주가 교실 수업에서 나타난 수행 양상을 바탕으로 구체화한 구성 요소이다.

FGI 분석에 따르면 학생들은 생성형 AI를 활용하여 문장을 생성하거나 변형하는 데에는 비교적 익숙하지만, AI 산출물을 수행 맥락 속에서 재검토해야 할 대상으로 인식하지 못하는 경우가 많았다. 특히 과제 목적이나 글의 흐름, 개인 경험과의 정합성을 기준으로 내용을 점검하기보다 표현의 유창성이나 문장의 완결성을 근거로 그대로 수용하는 경향이 반복적으로 나타났다. 이러한 양상은 AI 산출물이 판단의 대상이 아니라 곧바로 채택되는 결과물로 기능하게 만드는 요인으로 작용할 수 있다.

재맥락화가 충분히 이루어지지 않을 경우 이후 단계에서 요구되는 사실성 검증이나 책임 귀속 판단 역시 충분히 작동하기 어렵다. 이는 AI 활용 환경에서 언어 수행의 판단 과정이 자동적으로 이루어지지 않음을 보여주며, 학습자가 AI 산출물을 자신의 언어 활동 속에서 다시 해석하고 재구성하는 과정이 필요함을 시사한다.

또한 재맥락화 수행은 AI에 대한 인식 방식과도 관련된다. AI 산출물을 결과물이 아니라 수정과 재구성이 가능한 텍스트로 인식할 때 글의 구조 조정이나 표현 수정 시도가 상대적으로 활발하게 나타났다. 이러한 점에서 재

맥락화는 이후 상호작용을 통해 의미를 조정할 수 있는 판단의 출발점으로 이해할 수 있다.

이에 본 연구는 교실 맥락에서 관찰 가능한 수행 행동을 중심으로 재맥락화 수행 지표를 <표 6>과 같이 정리하였다.

<표 6> 재맥락화 수행 지표

수행 지표	관찰 가능한 수행 행동
AI 산출물을 초안으로 인식하기	AI 문장을 '제출용 결과물'과 '조정 가능한 초안'으로 구분하여 설명하고 수정이 필요한 이유를 과제 맥락에 근거해 진술함
구조 및 조직을 기준으로 선택·삭제·재배열하기	글의 목적을 기준으로 AI 문장 일부를 삭제·이동·재배열하고 그 판단 근거를 논리 구조 차원에서 설명함
맥락·청자 적합성을 기준으로 수용 여부 판단하기	표현의 자연스러움이 아니라 담화 목적·청자 특성을 기준으로 수용 또는 수정 여부를 결정하고 그 이유를 제시함
AI 응답을 조정 가능한 텍스트로 인식하기	AI 응답의 한계를 인식하고 이후 수정이나 질문 조정을 통해 보완할 수 있음을 설명함

2) 검증

검증은 생성형 AI가 제시한 응답을 그대로 수용하지 않고, 내용의 사실성, 근거의 타당성, 관점이나 편향의 가능성을 점검하는 판단 수행을 의미한다. 학습자는 AI 응답의 내용이 과제 목적에 적절한지 확인하고, 필요한 경우 추가 자료를 통해 정보의 정확성을 확인하거나 관점을 비교할 수 있다. 이러한 개념은 예비 범주 가운데 비판·검증 범주가 교실 맥락에서 나타난 수행 양상을 반영하여 구체화한 구성 요소이다.

FGI 분석에 따르면 학생들은 AI 응답의 유창성이나 문장의 완결성을 신체의 근거로 삼아 출처 확인이나 사실 점검을 수행하지 않는 경우가 많았다. 이러한 양상은 AI 환경에서 비판적 읽기 절차가 자동적으로 작동하지 않을 가능성을 보여준다.

이와 같은 상황에서는 AI 응답의 내용이 사실과 다르거나 특정 관점에

치우쳐 있더라도 그 문제가 충분히 드러나지 않을 수 있다. 특히 출처 확인이나 외부 자료와의 대조가 이루어지지 않을 경우 학습자는 자신의 언어 수행에서 정보의 신뢰성을 판단하는 경험을 충분히 축적하기 어렵다. 이러한 점에서 검증 수행은 재맥락화 이후 산출물의 타당성을 점검하는 핵심 판단 단계로 이해할 수 있다.

이에 본 연구는 교실 맥락에서 관찰 가능한 수행 행동을 중심으로 검증 수행 지표를 <표 7>과 같이 정리하였다.

<표 7> 검증 수행 지표

수행 지표	관찰 가능한 수행 행동
사실성과 정확성 점검하기	AI 응답의 핵심 주장 중 오류 가능성이 있는 내용을 지적하고, 수정 또는 보류의 이유를 명시함
출처 및 신뢰도 대조하기	AI 응답의 핵심 정보를 외부 자료와 대조하여 일치·불일치 여부를 확인하고 판단 근거를 기록함
관점 및 편향 읽어내기	AI 텍스트가 전제하는 관점이나 배제된 시각을 식별하고, 그 한계를 언어적으로 설명함

3) 책임 귀속

책임 귀속은 생성형 AI가 산출한 문장을 활용하더라도 표현의 선택·수정·제시 과정에 대한 최종 판단이 학습자에게 있다는 점을 인식하고, 그 판단 근거를 설명할 수 있는 수행을 의미한다. 생성형 AI 환경에서는 텍스트를 생성한 주체와 그것을 선택하여 사용하는 주체가 분리될 수 있으므로, 책임은 결과 자체보다 해당 표현을 선택하고 수정하여 제시한 판단 과정에 초점을 둔다. 이러한 개념은 예비 범주 가운데 윤리·책임 범주가 교실 수업에서 나타난 수행 양상을 바탕으로 구체화한 구성 요소이다.

FGI 분석에 따르면 오류나 부적절한 표현이 지적될 때 학생들이 “AI가 썼다”라는 설명을 통해 발화 책임을 분리하는 반응이 반복적으로 관찰되었다. 이러한 양상은 생성형 AI가 텍스트 생산 과정에 개입하는 상황에서 발화

주체성이 약화됨을 보여준다.

국어과 언어 활동에서는 표현을 직접 생성했는지 여부보다 해당 표현을 선택하고 제시한 주체가 누구인가가 책임 판단의 기준이 된다. 그러나 AI 환경에서는 생성 행위와 채택 행위가 분리되면서 발화 책임이 충분히 인식되지 않는 상황이 나타날 수 있다. 따라서 책임 귀속 수행은 표현 선택과 제시에 대한 최종 판단을 자신의 언어 행위로 인식하는 단계로 이해할 수 있다. 특히 표현의 타당성이나 윤리적 영향에 대한 고려 없이 AI 산출물을 그대로 제출할 경우 발화 주체성에 대한 인식이 약화될 수 있다.

이에 본 연구는 교실 맥락에서 관찰 가능한 수행 행동을 중심으로 책임 귀속 수행 지표를 <표 8>과 같이 정리하였다.

<표 8> 책임 귀속 수행 지표

수행 지표	관찰 가능한 수행 행동
발화 주체성 인식하기	AI가 생성한 문장을 선택·수정·제출한 이유를 자신의 판단으로 설명하며, 결과에 대한 책임을 명시함
출처 및 표절 가능성 판단하기	AI 산출물 사용 시 인용 여부를 결정하고, 그 판단 기준을 과제 요구와 윤리 기준에 근거해 제시함
사회적 영향 성찰하기	자신의 표현이 오해·왜곡·차별을 유발할 가능성을 언급하고, 수정 또는 유지의 이유를 설명함

4) 의미 협상

의미 협상은 생성형 AI를 단순히 답을 제공하는 도구로 사용하는 것이 아니라, 질문과 응답의 상호작용을 통해 내용을 점검하고 조정하는 수행 능력을 의미한다. 학습자는 AI의 응답을 그대로 수용하기보다 질문을 수정하거나 추가 질문을 제시하여 내용을 보완하고, 여러 응답을 비교하면서 과제 목적에 맞는 의미를 구성할 수 있다. 이러한 개념은 예비 범주 가운데 상호작용·협상 범주가 교실 맥락에서 나타난 수행 양상을 반영하여 도출된 구성 요소이다.

FGI 분석에 따르면 학생들의 생성형 AI 활용은 대체로 “한 번 묻고 한 번 답을 받는다” 단회적 구조에 머무르는 경우가 많았다. 이러한 상황에서는 AI 응답을 바탕으로 의미를 점진적으로 조정하거나 관점을 확장하는 과정이 충분히 이루어지지 않으며 사고가 초기 단계에서 종결되는 경향이 나타났다. 특히 AI 응답의 한계를 인식하더라도 질문의 조건이나 범위를 수정하여 재질문하는 시도는 상대적으로 드물게 나타났다.

이러한 양상은 AI와의 상호작용이 탐색적 대화 과정으로 조직되기보다 결과를 얻기 위한 도구적 요청으로 활용되는 경향과 관련된다. 그러나 생성형 AI 환경에서 언어 수행의 판단은 단회적 응답 수용이 아니라 응답의 한계를 인식하고 질문을 수정하거나 복수의 관점을 비교하는 과정 속에서 형성될 수 있다.

이에 본 연구는 교실 맥락에서 관찰 가능한 수행 행동을 중심으로 의미 협상 수행 지표를 <표 9>와 같이 정리하였다.

<표 9> 의미 협상 수행 지표

수행 지표	관찰 가능한 수행 행동
질문 조건 조정하기	초기 AI 응답의 한계를 인식하고 질문의 조건·범위·초점을 수정하여 재질문함
후속 질문 생성하기	AI 응답의 모호성이나 불충분한 설명을 근거로 추가 질문을 생성함
응답 비교 및 선별하기	복수의 AI 응답을 비교하여 과제 목적에 더 적합한 관점이나 근거를 선택함
의미 재구성하기	AI 응답을 자신의 경험·과제 맥락·답화 목적에 맞게 재조직하여 최종 발화로 구성함

3. 국어과 기반 AI 리터러시 요소의 관계와 수업 적용

1) 국어과 기반 AI 리터러시 요소의 순환적 판단 구조

이 연구에서 도출한 네 가지 국어과 기반 AI 리터러시 요소는 분리된 역

량의 병렬적 목록이 아니라, 생성형 AI 활용 과정에서 학습자의 언어적 판단이 반복적으로 조직되는 순환적 판단 구조로 이해할 수 있다. 각 요소는 고정된 선형 단계로 작동하기보다 실제 언어 수행 과정에서 상호 연계되며, 특정 판단 과정에서 드러난 문제가 다시 앞선 판단을 호출하는 방식으로 반복적으로 작동한다.

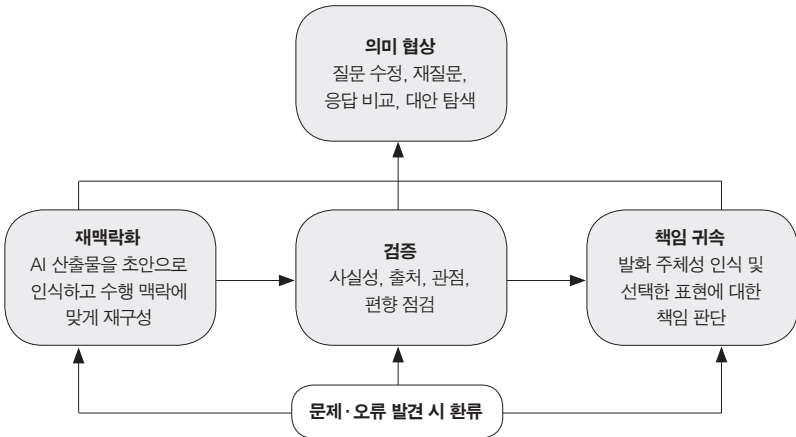
이러한 구조는 생성형 AI의 작동 방식과도 관련된다. 생성형 AI는 확률적 예측을 기반으로 텍스트를 생성하기 때문에, 산출물이 수행 맥락에 적합한 의미를 갖기 위해서는 인간의 검토와 수정, 가치 판단이 지속적으로 개입될 필요가 있다. 김연지·서혁(2025: 54-55)도 생성형 AI 산출물이 창의적으로 의미 있는 결과로 기능하기 위해서는, 인간이 문제 발견과 결과 평가의 단계에서 중심적으로 개입하여 가치와 의미를 부여하는 판단 과정이 중요함을 강조하였다.

이 연구가 제안하는 구조에서 생성형 AI 활용은 재맥락화 판단에서 출발한다. 학습자는 AI 산출물을 완성된 답이 아니라 과제 목적, 장르, 청자, 담화 맥락에 맞게 재구성해야 할 초안 텍스트로 인식하고 내용을 점검한다. 다음으로 검증 판단이 이루어지며, 학습자는 응답의 사실성, 출처, 논리적 타당성, 관점과 편향을 점검한다. 이 과정에서 오류가 발견될 경우 학습자는 다시 재맥락화를 수행하거나 새로운 정보를 탐색하게 된다.

이후 책임 귀속 판단이 작동한다. 학습자는 최종적으로 선택한 표현이 자신의 발화 행위에 포함된다는 점을 인식하고, AI 제안과 자신의 판단에 따른 선택을 구분하게 된다. 마지막으로 의미 협상 과정에서는 질문 수정, 재질문, 응답 비교, 대안 탐색과 같은 상호작용이 이루어지며 학습자는 의미를 점진적으로 조정하고 재구성한다. 이 과정에서 드러난 문제나 새로운 정보는 다시 재맥락화와 검증 판단을 호출하며 판단 과정은 반복된다.

이와 같이 재맥락화, 검증, 책임 귀속, 의미 협상의 네 요소는 일정한 판단 흐름을 형성하지만 실제 수행에서는 특정 단계에서 발견된 문제나 새로운 정보가 다시 앞선 판단을 호출하며 반복적으로 재구성된다. 따라서 국어

과 기반 AI 리터러시는 고정된 단계 모형이 아니라 언어 수행 과정에서 지속적으로 갱신되는 판단 체계로 이해할 수 있다. 본 연구는 이러한 관계를 <그림 1>과 같이 제시하였다.



<그림 1> 국어과 기반 AI 리터러시 요소의 순환적 의미 구성 구조

2) 국어과 기반 AI 리터러시를 반영한 수업 과제

이 연구는 생성형 AI 활용 과정에서 나타나는 학생들의 어려움을 기술 활용의 미숙함이 아니라, 국어과 언어 수행 과정에서 판단이 중단되거나 유보되는 구조적 문제로 해석하였다. 이러한 관점에서 도출된 재맥락화, 검증, 책임 귀속, 의미 협상의 네 요소는 추상적 역량 진술에 머무르기보다 실제 수업 설계와 평가 장면에서 관찰·지도·피드백이 가능한 수행 단위로 구체화될 필요가 있다.

이에 본 연구는 국어과 기반 AI 리터러시를 반영한 수업 과제 예시를 제시하여, 제한한 개념 틀이 교실 맥락에서 어떻게 적용될 수 있는지 살펴보고자 한다. 이는 AI 리터러시 논의를 국어과 수업의 과제 구조와 평가 준거와 연결하여 검토하기 위한 하나의 사례로 이해될 수 있다. 과제 유형과 수업 맥락, 과제 구성은 <표 10>에 제시하였다.

〈표 10〉 국어과 기반 AI 리터러시를 반영한 수업 과제(예시)

- 과제 유형: 토론 준비를 위한 AI 활용 자료 구성 과제
- 수업 맥락: 고등학교 화법 수업 / 쟁점 토론 사전 준비 단계

요소	과제 조건	산출물 형태	평가 준거
재맥락화	AI가 생성한 토론 자료를 그대로 사용하지 말고, 토론 주제·입장·청자 조건에 맞게 재구성할 것	AI 원문과 수정본 비교표	AI 산출물을 초안으로 인식하고, 목적·담화 맥락에 따라 선택·삭제·재배열했는가
검증	AI 응답 중 핵심 주장 1~2개를 선정하여 사실성·출처를 외부 자료로 대조할 것	출처 대조표, 오류·보완 메모	사실성·출처·관점을 기준으로 비판적 점검이 이루어졌는가
책임 귀속	최종 토론 자료에서 AI 사용 여부와 선택·수정의 이유를 명시할 것	선택 근거 진술문	표현 선택의 책임이 학습자에게 귀속된다는 인식이 언어적으로 드러나는가
의미 협상	초기 AI 응답 이후, 최소 1회 이상 질문을 수정·보완하여 재질의할 것	질문 수정 로그, 응답 비교 기록	질문 조정과 응답 비교를 통해 의미가 확장·조정되었는가

〈표 10〉에서 중요한 점은 AI 활용 여부 자체가 아니라, AI 활용 과정에서 학습자의 판단이 언제 어떤 기준으로 개입했는가가 산출물에 드러나도록 과제를 설계했다는 점이다. 수정 로그, 비교표, 선택 근거 진술과 같은 산출물은 언어 수행 결과뿐 아니라 그 과정에서 작동한 판단 구조를 함께 드러낸다. 이러한 과제 설계는 말하기 수행평가에서 종종 간과되어 온 ‘왜 그렇게 말했는가’, ‘왜 이 표현을 선택했는가’와 같은 판단의 근거를 평가 장면으로 끌어들이는 효과를 지닌다.

특히 재맥락화-검증-책임 귀속-의미 협상으로 이어지는 판단 구조는 토론, 발표, 협상, 글쓰기 등 다양한 국어과 수행 과제로 확장 적용될 수 있다.

다만 본 연구에서 제시한 과제 예시는 하나의 적용 가능성을 보여주는 사례이며, 실제 교실에서는 학습자의 수준, 수업 시간, 평가 목적에 따라 각 요소의 강조점과 과제 구조가 조정될 필요가 있다.

V. 결론 및 교육적 시사점

이 연구는 생성형 인공지능 활용이 일상화된 국어 교실에서 관찰되는 학습자의 어려움을 기술적 숙련의 문제로 환원하기보다, 언어적 판단 구조가 실제 수업 맥락에서 어떻게 작동하고 있는지를 국어과의 관점에서 탐색하고자 하였다. 이를 위해 국어 교사 7인을 대상으로 포커스 그룹 인터뷰와 문헌 분석을 실시하여, 생성형 AI 활용 과정에서 반복적으로 나타나는 문제양상을 분석하고 이에 대응하는 국어과 기반 AI 리터러시의 구성 요소와 통합적 구조를 탐색적으로 도출하였다.

연구 결과, 교사들이 인식한 학생들의 어려움은 생성형 AI 기능 자체에 대한 미숙함이라기보다, AI 산출물을 완성된 답안으로 전제하는 인식에서 출발하여 목적과 장르, 맥락에 따른 재구성이 이루어지지 못하고, 그 결과 검증의 생략과 책임의 전가, 상호작용의 단절로 연쇄적으로 확장되는 양상으로 나타났다. 이는 생성형 AI 활용의 문제가 개별 기능의 부족이 아니라, 국어과 언어 활동에서 핵심적으로 작동해야 할 재맥락화, 검증, 책임 귀속, 의미 협상의 판단 과정이 충분히 조직되지 못한 상태와 밀접하게 관련됨을 시사한다.

이에 이 연구는 국어과 기반 AI 리터러시를 AI를 잘 사용하는 능력이 아니라, ‘AI와 함께 수행되는 언어 활동에서 판단과 책임을 회복하는 능력’으로 정의하고, 이를 재맥락화, 검증, 책임 귀속, 의미 협상의 네 요소로 정리하였다. 이들 요소는 분리된 기능의 나열이 아니라, 생성형 AI 활용 과정에서 학습자가 의미 구성의 주체로 반복적으로 개입하도록 돕는 순환적 판단 구조로 상호 연계되어 작동한다. 즉 재맥락화 판단을 출발점으로 검증과 책임 귀속 판단이 이루어지고, 그 결과가 의미 협상을 통해 조정되며, 이러한 판단 과정이 수행 전반에서 반복되는 구조로 이해할 수 있다.

이러한 연구 결과는 생성형 AI 활용 지도가 기능적 사용법이나 프롬프

트 기술의 습득에 머물러서는 안 되며, AI 산출물을 판단의 대상으로 인식하게 하는 재맥락화를 수업의 출발점으로 설정할 필요가 있음을 시사한다. 또한 비판적 읽기와 언어 윤리는 정보 확인이나 규범 전달 차원에 머무르기보다, AI 텍스트의 사실성과 관점, 책임 귀속을 점검하고 질문 조정과 재질의를 통해 의미를 협상하는 판단 과정 중심의 수업 설계로 통합될 필요가 있다.

종합하면 이 연구는 생성형 AI 활용 문제를 기술적 도구 활용이 아니라 언어 수행에서의 판단 구조의 문제로 살펴보았다는 점에서 의의를 지닌다. 이는 생성형 AI 환경에서 국어과 교육이 단순한 활용 능력보다 의미를 판단하고 조정하며 책임 있게 표현하는 언어적 주체성의 형성을 중요하게 다루어야 함을 시사한다. 한편 본 연구는 질적 탐색 연구로 참여 교사 수와 연구 맥락에 제한이 있으며 학습자 수행 자료를 직접 분석하지 못한 한계를 지닌다. 향후 연구에서는 제안한 국어과 기반 AI 리터러시 요소를 바탕으로 수업 설계와 평가 도구를 개발하고, 학습자의 발화·쓰기·성찰 자료 분석을 통해 추가적인 검토가 이루어질 필요가 있다.

* 본 논문은 2026.01.08. 투고되었으며, 2026.02.08. 심사가 시작되어 2026.03.07. 심사가 종료되었음.

참고문헌

- 김연지·서혁(2025), 「생성형 AI 시대의 창의성과 프롬프트 리터러시」, 『국어교육학연구』 60(1), 41-82.
- 김운용(2025), 「감응력 없는 저자 - 생성형 AI 시대의 리터러시와 비판적 사고」, 『리터러시 연구』 16(5), 389-416.
- 김인주·이소율·김귀훈(2024), 「초·중등교사의 생성형 AI 리터러시 측정도구 개발 및 타당화 연구」, 『학습자중심교과교육연구』 24(22), 697-708.
- 김종규(2023), 「생성형 인공지능, 생각하는 존재(homo cogitans) 그리고 리터러시 교육의 방향」, 『사고와표현』 16(2), 7-31.
- 김태윤·문수연·백지혜·이소정(2025), 「디지털 리터러시 역량 강화를 위한 데이터 기반 AI·과학·기술 융합 수업모형 개발」, 『학습자중심교과교육연구』 25(4), 987-1006.
- 노환호·김현정·김민진(2024), 「한국형 생성 인공지능 리터러시 척도 개발 및 타당화」, 『지식경영연구』 25(3), 145-171.
- 성은모·김동호·신서경·이영주(2023), 「인공지능(AI)의 교육적 실천을 위한 가능성과 과제」, 『교육공학연구』 39(4), 1479-1508.
- 이유미(2021), 「AI 시대의 리터러시 특성에 관한 연구 - AI 리터러시와 관계 리터러시를 중심으로」, 『어문연구』 110, 281-302.
- 이유미·박윤수(2021), 「AI 리터러시 개념 설정과 교양교육 설계를 위한 연구」, 『어문론집』 85, 451-474.
- 이지영(2024), 「AI를 활용한 리터러시 교육 관련 국내 연구와 언론 기사에서 나타난 키워드 비교 연구」, 『작문연구』 60, 97-141.
- 정서현·박주연(2024), 「AI 리터러시 역량 결정 요인 연구: AI 이용 경험과 혁신성을 중심으로」, 『방송통신연구』 128, 137-168.
- 정현선(2004), 「디지털 리터러시의 국어교육적 고찰」, 『국어교육학연구』 21, 5-42.
- 최숙영(2022), 「AI 리터러시 프레임워크에 대한 연구」, 『컴퓨터교육학회 논문지』 25(5), 73-84.
- 최재호·전신영(2023), 「인공지능 시대의 디지털 역량 함양을 위한 디지털 리터러시 역량 평가 도구 개발 및 타당화」, 『대학 교수-학습 연구』 16(3), 95-122.
- 최진영(2023), 「우리나라 중·고등학생의 디지털 미디어 리터러시 인식 연구」, 『리터러시 연구』 14(2), 241-276.
- 최창원(2023), 「초등 국어과 AI 리터러시 융합 교육을 위한 예비적 고찰」, 『한국초등국어교육』 76, 434-458.
- 황현정·황용석·박지수·신민호·이현중(2024), 「AI 리터러시 측정 도구 개발과 타당화 연구」, 『리터러시 연구』 15(2), 247-278.
- Wang, B., Rau, P. L. P., & Yuan, T. (2023), "Measuring user competence in using artificial intelligence: Validity and reliability of artificial intelligence literacy scale", *Behaviour & Information Technology*, 42(9), 1324-1337.

생성형 AI 활용 국어 수업의 문제 양상과 국어과 기반 AI 리터러시 구성 요소 — 국어 교사 FGI 질적 연구

허선아

이 연구는 생성형 인공지능 활용이 보편화된 국어 교실에서 나타나는 학습자의 어려움을 기술 숙련의 문제로만 보기보다 국어과 언어 활동에서의 판단 과정의 관점에서 살펴보고, 이에 기반한 국어과 AI 리터러시의 구성 요소를 탐색하고자 하였다. 이를 위해 국어 교사 7인을 대상으로 포커스 그룹 인터뷰(FGI)와 문헌 분석을 결합한 질적 탐색 연구를 수행하였다. 분석 결과, 교사들은 학생들이 생성형 AI 산출물을 완성된 답안으로 수용하고 과제 맥락에 맞게 재구성하거나 검증하는 과정, 표현 선택에 대한 책임 인식, 질문과 재질문을 통한 의미 조정이 충분히 이루어지지 않는 양상을 공통적으로 인식하고 있었다. 이는 생성형 AI 활용 문제를 개별 기술의 부족이라기보다 국어과 언어 수행에서 판단 과정이 충분히 조직되지 못한 상태와 관련된 현상으로 이해할 필요가 있음을 시사한다. 이에 이 연구에서는 국어과 기반 AI 리터러시를 재맥락화, 검증, 책임 귀속, 의미 협상의 네 구성 요소로 구조화하였다.

핵심어 생성형 인공지능, 국어 교육, AI 리터러시, 언어적 판단, 교사 인식, 질적 연구

ABSTRACT

AI Use in Korean Language Classrooms: Teachers' Perspectives and Components of Korean Language-Based AI Literacy

Heo Seona

This study examines the challenges encountered in Korean language classrooms employing generative AI-framing them as issues related to language-based judgment rather than technical proficiency. Based on a literature review and focus group interviews with seven Korean language teachers-the results indicated that students tend to treat AI-generated texts as completed answers, with limited recontextualization, verification, responsibility attribution, and interaction.

Accordingly, the study proposes four components of Korean language-based AI literacy: recontextualization, verification, responsibility attribution, and meaning negotiation. Therefore, this contributes to a subject-specific AI literacy framework that emphasizes judgment and responsibility in Korean language education.

KEYWORDS Generative AI, Korean language education, AI literacy, language-based judgment, teacher perception, qualitative study