

인공지능(AI · LLM)의 전문직 대체 가능성과 사회적 압력 기반 규제샌드박스 설계에 관한 연구

유 한 별* · 김 승 완**

본 연구는 생성형 AI의 발전이 의사, 변호사 등 전문직종에 미치는 영향을 기술적 대체 가능성과 사회적 압력이라는 이원적 관점에서 분석하고, 이에 적합한 한국형 규제샌드박스 운영 모형을 제안하는 것을 목적으로 한다. 2022년부터 2026년까지의 중앙일간지 뉴스 자료를 통해 사회적 압력을 확인하고, LLM(GPT·Claude·Gemini)을 활용한 업무 대체율 평가 데이터를 결합하여 14개 전문직종의 기술적 대체가능성을 평가하고, 사회적 압력 자료가 충분한 직종을 중심으로 위험도를 산출하였다. 단일 모델 의존을 피하기 위해 3종 모델의 평가를 앙상블하고, 모델 간 측정 신뢰도를 ICC·Krippendorff α ·Kendall W로 검증하였다. 분석 결과, 법무사·관세사는 높은 기술적 대체율을 보이거나 상대적으로 낮은 사회적 주목도를 보인 반면, 의사·변호사는 기술적 대체율은 중간 수준이나 사회적 압력이 극도로 높은 고갈등 영역으로 나타났다. 이에 본 연구는 포획 이론과 반응적 규제 이론을 바탕으로, 전문직종의 특성을 고려한 4단계 규제샌드박스 모형을 설

* 주저자, 선문대학교 행정·공기업학과, 디지털콘텐츠학과 조교수, 충남 아산시 탕정면 선문로 221번길 70 (hanbyeolyoo@sunmoon.ac.kr)

** 교신저자, 한경국립대학교 사회통합학부 공공행정전공 부교수, 경기도 평택시 한경대학교로 35(waniss33@hknu.ac.kr)

접수일: 2026.03.24. / 심사일: 2026.04.16. / 게재확정일: 2026.06.24.

계하여 특히 고위험 영역에서의 인간 개입 의무화와 저위험 행정 업무의 자동화 패스트 트랙을 구분하는 차등적 규제 전략을 제시하며, 2026년 시행된 한국 AI 기본법의 하위 법령 고도화 설계에 대한 정책적 시사점을 도출하였다.

핵심용어: 인공지능, 전문직, 표준화, 확률화, 규제샌드박스, 사회적 압력, 언론 프레임

I. 서론

인공지능(Artificial Intelligence, 이하 AI), 특히 거대언어모델(Large Language Model, 이하 LLM)의 급격한 발전은 지식 노동의 성격을 근본적으로 변화시키고 있다. 과거의 자동화가 육체노동이나 단순 반복 업무를 대상으로 했다면, 현재의 생성형 AI는 고도의 인지 능력이 요구되는 전문직 영역(법률, 의료, 회계 등)으로 그 영향력을 확장하고 있다. 전문직 자격증은 전통적으로 정보비대칭성을 근거로 시장 진입 장벽을 형성하고 서비스 품질을 보장하는 역할을 해왔으나, AI 기술의 보편화는 이러한 전문직의 해자(Moat)를 기술적으로 무력화할 가능성을 제기한다.

한국 사회에서 전문직은 국가가 공인한 자격 시험을 통과한 소수에게 업무 독점권을 부여하는 방식으로 관리되어왔다. 그러나 AI 기술의 발전은 이러한 독점적 업무 영역의 상당 부분을 기술적으로 자동화하거나 보조할 수 있는 가능성을 열었다. 특히 2026년 발효한 「인공지능 발전과 신뢰 기반 조성 등에 관한 기본법(이하 AI 기본법)」은 AI 기술의 안전한 활용을 위한 법적 토대를 마련하고 있으며, 이에 따라 각 전문직역별 규제 체계의 재편이 불가피한 상황이다.

다시 말해, 2024년 EU AI Act의 발효와 2026년 한국 AI 기본법 시행 등 전 세계적으로 AI 규제 거버넌스가 구체화되는 시점에서, 전문직역과 AI의 공존 모델을 모색하는 것은 시급한 정책 과제이다. 기존의 획일적인 규제 접근은 혁신을 저해하거나, 반대로 고위험 영역에서의 안전 공백을 초래할 수 있다. 특히 전문직 단체의 강력한 이익 집단화 경향과 기술 도입에 따른 사회적 갈등(예: 의료 대란, 법률 플랫폼 갈등)이 혼재된 한국적 상황에서는 기술적 대체 가능성뿐만 아니라 사회적 수용성을 고려한 정교한 규제 설계가 요구되는 시점이다.

이에 따라, 본 연구의 목적은 전문직종별 AI 대체 가능성과 사회적 압력을 실증적으로 분석하고, 이를 바탕으로 전문직 특화 AI 규제샌드박스 운영 모형을 설계하는 것이다. 이에 따라 연구질문은 다음과 같이 요약할 수 있다.

- (1) 전문직 업무에 대한 사회적 압력은 어떠한 수준으로 형성되어 있는가?
- (2) 전문직 업무의 표준화 수준과 확률화 수준은 직종별로 어떻게 다르며, 이에 따른 AI 대체율은 어떻게 산출되는가? 산출된 AI 대체율(기술적 요인)과 뉴스 미디어를 통해 표출된 사회적 압력(사회적 요인) 사이에는 어떤 관계가 존재하는가?
- (3) 기술적 대체 가능성과 사회적 압력의 상호작용을 고려할 때, 전문직종별로 차별화된 규제샌드박스 설계 원칙을 구체화 할 수 있는가?

구체적으로 첫째, 뉴스 텍스트 마이닝을 통해 각 직종에 가해지는 사회적 변화 압력을 정량화한다. 둘째, 생성형 AI, 즉, LLM(GPT, Claude, Gemini)을 활용하여 직종별 법령 기반 업무의 기술적 대체율을 측정한다. 셋째, 이 두 가지 변수를 축으로 전문직종의 위험도를 분류하고, 각 유형에 적합한 규제샌드박스 진입 및 기준을 제시한다. 또한 본 연구는 이론적으로 기존의 자동화 연구가 간과했던 표준화와 확률화라는 이원적 분석 프레임워크를 제시함으로써 AI와 전문직의 상호작용을 보다 정밀하게 설명한다. 방법론적으로는 LLM(GPT 5.2, Claude Opus 4.5, Gemini 3 등)을 활용한 전문직종 업무 평가와 뉴스 기사 텍스트마이닝을 결합한 혼합 연구방법을 시도했다는 점에 의의가 있다. 실무적으로는 AI 기본법 시행에 맞춰 정부가 하위 법령을 설계하고 전문직 관련 규제를 개혁하는 데 필요한 정책 근거를 마련하고, 향후 AI 기본법 시행령 고도화 개정 및 직역별 갈등 조정 메커니즘 구축에 기여할 수 있다고 판단된다.

II. 이론적 논의

1. 인공지능(AI) 및 전문직의 특성과 대체 가능성에 대한 전제조건

현대의 기계학습은 사람이 규칙을 일일이 설계하는 방식보다, 대규모 데이터에서 통계적 규칙(패턴)을 학습하여 예측·분류·생성 과제를 수행하는 데 비교우위를 가진다

(Brynjolfsson & Mitchell, 2017). 특히 LLM은 방대한 텍스트 데이터로부터 일반적인 언어 패턴을 학습해 다양한 과제를 수행할 수 있지만, 그 작동은 본질적으로 데이터 분포와 학습 신호(정답·피드백)에 의존하며, 업무 목표가 학습과 평가 가능한 형태로 주어질 때 성능 향상이 가능하다. 따라서 전문직 업무에서 AI의 도입 가능성은 ‘AI가 똑똑해졌는가’보다, ‘업무가 과업 단위로 분해되고 기계가 다룰 수 있는 입력·출력 구조를 갖추는가’라는 조건에 의해 설명 가능하다(Acemoglu & Restrepo, 2019).

전문직은 단순 직무기술의 집합이 아니라, 특정 문제 영역에 대한 관할권(jurisdiction)을 바탕으로 전문지식을 독점·관리하고 사회적 신뢰를 획득해 온 제도적 직업군이다(Abbott, 2014). 전문직 업무는 표준 지식만으로 환원되지 않는 재량(discretion)과 상황판단, 암묵지, 윤리적 의무(신뢰·비밀·충실의무 등) 및 책임(민형사·행정상 책임)과 결합되어 수행되는 경향이 크다(Evetts, 2003). 이 때문에 전문직의 AI 대체가능성은 단순히 ‘업무가 존재하는가’가 아니라, 전문직 업무가 가진 재량·공익성·책임성을 손상시키지 않으면서도 기계가 처리 가능한 형태로 변환될 수 있는지를 중심으로 검토되어야 한다(Evetts, 2003).

최근 AI가 어떤 전문직을 얼마나 대체할 수 있는지에 대해서 그 논의가 지속적으로 확산되고 있다¹⁾. AI의 전문직 대체가능성은 기술의 성능 자체보다, 직무가 AI가 다루기 쉬운 형태로 변환될 수 있는지(업무의 구조적 조건)에 의해 크게 좌우된다(Autor, Levy & Murnane, 2003; Brynjolfsson, Mitchell & Rock, 2018). LLM(기계학습·딥러닝·대규모 언어모델)은 규칙을 사람이 직접 작성하는 방식보다 대규모 데이터에서 패턴을 학습해 예측·분류·생성하는 방식에 강점을 가진다(Brown et al., 2020; Brynjolfsson & Mitchell, 2017). 따라서 AI가 전문직 업무에 실질적으로 투입되기 위해서는 최소한 다음의 두 조건이 충족될 필요가 있다(Acemoglu & Restrepo, 2019; Brynjolfsson & Mitchell, 2017).

- (1) (표준화) 자료 입력과 절차가 정형화되어 기계가 처리 가능한 형태로 정리
표준화는 특정 업무가 규칙·절차·서식·체크리스트로 코드화되어, 수행자가 달라도 결

1) “이제는 전문직도? 회계사·변호사도 신입 대신 AI 쓴다”, KBS 뉴스, 2026.01.27.,
https://news.kbs.co.kr/news/pc/view/view.do?ncd=8469949&req_by=aipick

과가 유사하게 재현될 수 있는 정도를 의미한다(Autor, Levy & Murnane, 2003). 표준화가 높을수록 업무는 명시적 규칙(ex ante rule)으로 표현되며, 규칙 기반 자동화(RPA, 규칙엔진, 문서 자동화 등)의 대상이 되기 쉽다(Autor, Levy & Murnane, 2003; Leshob, Bédard & Mili, 2020; Wewerka & Reichert, 2020). 즉, 표준화 처리는 명시적인 규칙(Explicit Rules)과 알고리즘에 기반하여 업무를 수행하는 AI의 능력을 의미한다. 이는 전통적인 전문가 시스템(Expert Systems)이나 로봇 프로세스 자동화(RPA) 기술이 대표적이다. 업무의 입력값과 출력값 사이에 명확한 'IF-THEN' 논리 구조가 존재할 때, AI는 인간보다 빠르고 정확하게 업무를 처리할 수 있다.

예를 들어, 세무사의 업무 중 과세표준에 따른 세액 계산이나 법무사의 등기 신청 서류 작성은 법령에 정해진 엄격한 서식과 절차를 따르므로 표준화 가능성이 매우 높다. 이러한 영역에서 AI는 인간의 실수를 줄이고 효율성을 극대화하는 도구로 작동한다. 그러나 규칙이 모호하거나 예외 상황이 빈번한 경우, 표준화 기반 AI의 성능은 급격히 저하된다는 한계가 있다.

(2) (확률화) 목표가 데이터 기반 학습 문제로 정의되어, 성능을 평가·개선할 수 있음

확률화는 업무의 판단·추론·분류가 데이터 기반 예측 문제로 환원되어, AI가 학습을 통해 성능을 개선할 수 있는 정도를 의미한다(Brynjolfsson & Mitchell, 2017; Brynjolfsson, Mitchell & Rock, 2018). 여기서 확률화는 AI가 확률적으로 판단한다는 기술적 서술이 아니라, 업무가 데이터로 관측되고, 학습 목표(정답/성과)가 정의되며, 피드백(오류·성과)이 축적되는 구조인이라는 학습가능성의 개념이다(Brynjolfsson & Mitchell, 2017; Brynjolfsson, Mitchell & Rock, 2018). 확률화 처리는 대량의 데이터를 학습하여 통계적 패턴을 인식하고, 이를 바탕으로 가장 가능성 높은 결과를 추론하거나 예측하는 AI의 능력을 말한다. 이는 딥러닝과 최근의 대규모 언어 모델(LLM)이 여기에 해당한다. 명시적인 규칙으로 정의하기 어려운 비정형 데이터(이미지, 자연어 등)를 처리하는 데에도 탁월하다.

예를 들어, 의료 영상 판독에서 AI가 수만 장의 X-ray 이미지를 학습하여 환부를 찾아내거나, 변호사가 수천 건의 판례를 분석하여 승소 확률을 예측하는 것이 대표적인 예다. 확률화에 기반한 AI는 정답이 아닌 최적의 근사값을 제시하며, 이는 전문직의 직관이

나 경험적 지식을 대체하는 강력한 기제가 된다. 그러나 확률적 작동 방식은 설명 가능성의 부재, 데이터 편향에 따른 차별 위험 등의 문제를 내포한다.

종합하면 상기한 표준화·확률화 틀은 Autor et al.(2003)의 ‘Routine/Non-routine’ 구분과 Brynjolfsson·Mitchell(2017)의 기계학습 적합성 개념을 전문직 과업 평가에 맞게 조작화·통합한 측정틀이다. 본 연구의 이론적 차별성은 이 측정틀 자체보다, 기술적 대체가능성과 사회적 압력을 결합한 2×2 위험 유형화 및 이를 규제수단 포트폴리오·샌드박스 운영규칙으로 구체화한 점에 있다.

2. 언론 프레임링과 사회적 압력

기술 도입 과정에서 대중의 인식은 규제 정책 결정에 결정적인 영향을 미친다. Goffman의 프레임 이론(Erving Goffman's Framing Theory)에 따르면, 미디어는 특정 이슈를 선택하고 강조함으로써 대중의 해석 틀을 형성한다(Ytreberg, 2002). 전문직 AI 대체와 관련하여 뉴스·미디어는 일자리 위협(Job Displacement) 프레임이나 서비스 혁신(Innovation) 프레임을 통해 사회적 압력을 형성한다. 본 연구에서는 뉴스 데이터의 양(Volume)과 논조(Tone)를 통해 이러한 사회적 압력을 측정하고, 이를 기술적 요인과 함께 규제 설계의 핵심 변수로 활용한다.

상세히 설명하면, 기술 도입과 규제 설계는 기술의 객관적 위험만으로 결정되지 않고, 그 위험과 편익이 사회적으로 어떻게 해석·평가되는지에 의해 크게 좌우된다. 특히 신기술 분야에서는 제도 설계의 정당성(legitimacy)과 수용성을 확보하기 위해 이해관계자·대중의 인식이 중요한 정책 입력으로 작동하며, 규제샌드박스 역시 현장 참여자의 인식·정책수요가 제도 운영·개선에서 핵심 쟁점이 된다(최해옥·이광호, 2021).

대중 인식이 형성되는 메커니즘을 설명하는 대표 틀이 어빙 고프먼의 프레임 관점이며, ‘어떤 측면을 선택적으로 부각(salience)하느냐’가 해석과 판단을 구조화한다는 프레임링은 단순 정보 전달이 아니라 문제 정의·인과 귀인·도덕적 평가·해결책 제안의 방향을 함께 설정하며, 실험·메타 수준의 논의에서도 프레임이 공중의 태도와 정책 선호에 유의미한 영향을 미친다는 점이 반복 확인된다(Entman, 1993).

전문직 AI 대체 이슈에서는 미디어가 주로 일자리 위협(Job Displacement)과 서비

스 혁신(Innovation) 같은 경쟁적 프레임을 구성하면서 사회적 압력의 방향을 달리 만든다. 국내 선행연구에서도 AI 관련 보도에서 내용·방향 프레임이 체계적으로 분화됨이 보고되고, AI 보도 프레임이 수용자의 정서(희망/분노 등)와 행동 의향으로 연결된다는 분석·실험 결과가 제시된다(최민음·조은수, 2024). 또한 국가·언론 환경에 따라 AI 프레임이 규제·윤리·거버넌스 쟁점과 생산성·거시효과 등을 다르게 강조한다는 비교 연구는, 같은 주제이더라도 사회적 압력이 서로 다른 규제 경로를 촉발할 수 있음을 시사한다(박주현 외, 2024).

따라서 본 연구에서는 사회적 압력을 뉴스 자료의 수(Volume)와 논조(Tone)로 조작적 정의로 사용하고, 기사 수는 이슈 주목도 및 의제 강도의 근사치로 활용될 수 있으며 국내에서도 한국언론진흥재단의 빅카인즈(BigKinds) 자료를 통해 기사량 변화와 이슈 구조를 파악한 연구가 다수 존재한다(권충훈, 2019; 김경동 외, 2020; 도홍일 외, 2024; 박주현 외, 2024; 오준협·고보선, 2022; 이용희, 2020 등). 나아가 논조(긍·부정, 위협/기대) 지표는 뉴스 소비자의 관심·평가와 연결될 수 있다는 근거가 제시되어, 기술적 대체요인(표준화·확률화)과 사회적 압력(Volume·Tone)을 결합해 직종별 규제 강도와 샌드박스 조건(진입요건-운영조건-측정-종료 및 전환)을 차등 설계하는 분석틀을 강화할 수 있다고 판단되므로, 뉴스 기사 빈도 수 분석과 텍스트마이닝(워드클라우드, 네트워크분석, 논조 파악을 위한 감성분석)을 진행하기로 한다.

3. 전문직에 대한 규제 이론적 연계와 AI 관련 규제에 관한 국제 동향

전문직 규제 이론은 공익 이론과 포획 이론으로 구별된다. Posner(1974)의 공익 이론에 따르면 전문직 규제는 시장 실패를 교정하고 무자격자에 의한 소비자 피해를 방지하기 위해 존재하며, AI 맥락에서는 오작동·편향으로부터 소비자를 보호하는 논리로 이어진다. 특히 기계학습 기반 판단은 모델·데이터·절차의 복잡성으로 결과 산출 근거가 충분히 설명되지 않아 책임소재 규명과 사후구제를 어렵게 한다(Burrell, 2016).

반면 Stigler(1971)의 포획 이론은 규제기관이 피규제 집단(전문직 협회)의 이익에 포획되어 규제가 진입장벽 강화 수단으로 변질될 수 있음을 경고한다. ‘타다금지법’·‘로톡’ 사태처럼 직역 단체가 혁신 기술 도입을 저지하기 위해 규제를 활용하는 현상이 그 예이

다. Ayres & Braithwaite(1992)의 반응적 규제 이론은 피규제자 행태에 따라 자율규제에서 강제명령으로 강도를 조절하는 규제 피라미드를 제시하며, 이는 불확실성에 유연하게 대응하는 샌드박스 제도의 이론적 토대가 된다(권용수, 2022; 박종준, 2020).

AI 규제 국제 동향은 본 연구와 직접 관련된 위험기반 접근과 샌드박스 조항을 중심으로 간결히 정리한다. EU AI Act(Regulation (EU) 2024/1689)는 2024년 8월 발효되어 AI 시스템을 ①수용 불가능한 위험 ②고위험 ③제한적 위험 ④최소 위험의 4단계로 분류하고 차등 의무를 부과한다(〈표 1〉). 특히 교육·고용·필수서비스(의료·법률 등) 접근에 영향을 미치는 AI는 고위험으로 분류되며, 제57~59조는 회원국이 2026년 8월까지 최소 하나의 AI 규제샌드박스를 ‘통제된 환경(controlled environment)’으로 설립하도록 한다. 영국 FCA(2016~)·싱가포르 MAS(2016~)의 핀테크 샌드박스 역시 고위험 기술이라도 통제된 환경에서의 실증을 허용해 혁신과 안전의 균형을 모색한 사례로, 샌드박스 진입 기업의 자본조달·존속 가능성이 유의하게 개선되었음이 보고된다(Cornelli et al., 2024).

〈표 1〉 EU AI Act 위험등급 및 규제내용

위험 등급	규제 내용	법적 근거
수용 불가 위험	전면 금지 (사회적 신용평가, 실시간 원격 생체인식, 행동조작 AI 등 8가지 관행)	Article 5
고위험	엄격한 적합성 평가, 기술문서화, 위험관리체계, 인적 감독 의무 부과	Article 6, Annex III
제한적 위험	투명성 의무 (AI 사용 사실 고지, 딥페이크 표시 등)	Article 50
최소 위험	규제 없음 (AI 게임, 스팸 필터 등 대부분의 AI 시스템)	-

국내 「AI 기본법」(법률 제20676호, 2025.1.21. 제정, 2026. 1. 22. 시행)은 제2조 제4호에서 고영향 인공지능을 ‘사람의 생명·신체의 안전 및 기본권에 중대한 영향을 미치거나 위험을 초래할 우려가 있는 인공지능’으로 정의하고 보건의료·교통·교육 등 영역을 예시한다. 인공지능기본법 시행령은 사용영역·기본권 위협의 영향·중대성·빈도 등을 종합 고려하도록 하여 EU AI Act Annex III와 유사한 영역별 접근을 취하되 보다 유연한 판단 기준을 채택한다. 즉 한국의 AI 기본법 역시 위험기반 접근과 샌드박스 조항을 포함하므로, 본 연구는 AI의 전문직 대체 가능성·사회적 관심도·규제 이론·샌드박스 개

념을 종합하여 정책 함의를 도출하고자 한다.

III. 연구방법

1. 전문직 정의 및 자료와 분석방법

본 연구에서의 전문직의 정의는 일반적으로 사회 통념상 '8대 전문직'²⁾으로 불리는 직종과 보건복지부 주요 면허인 '의료인 및 의료기사'³⁾로 정의한다. 해당 전문직과 관련된 법령부터 소관부처는 아래 <표 2>와 같다.

<표 2> 전문직의 정의 및 법령, 소관부처

분류	전문직	관련 법령	시행령	시행규칙	협회명	소관부처
8대 전문직	변호사	변호사법	변호사법 시행령	변호사시험법 시행규칙	대한변호사협회	법무부
	변리사	변리사법	변리사법 시행령	변리사법 시행규칙	대한변리사회	특허청
	공인회계사	공인회계사법	공인회계사법 시행령	공인회계사법 시행규칙	한국공인회계사회	금융위원회
	세무사	세무사법	세무사법 시행령	세무사법 시행규칙	한국세무사회	기획재정부
	법무사	법무사법	-	법무사규칙	대한법무사협회	법무부
	공인노무사	공인노무사법	공인노무사법 시행령	공인노무사법 시행규칙	한국공인노무사회	고용노동부
	관세사	관세사법	관세사법 시행령	관세사법 시행규칙	한국관세사회	관세청
	감정평가사	감정평가 및 감정평가사에 관한 법률	감정평가 및 감정평가사에 관한 법률 시행령	감정평가 및 감정평가사에 관한 법률 시행규칙	한국감정평가협회	국토교통부

2) '8대 전문직' 불리는 문과생·작년 8만명 지원 '역대급', 한국경제, 2024.03.03.,

<https://v.daum.net/v/xawROfkPTu>

3) 의사, 치과의사, 한의사, 약사, 간호사, 의료기사

보건 의료 전문직	의사	의료법	의료법 시행령	의료법 시행규칙	대한 의사협회	보건 복지부
	치과 의사	의료법	의료법 시행령	의료법 시행규칙	대한치과 의사협회	보건 복지부
	한의사	의료법	의료법 시행령	의료법 시행규칙	대한한의사협회	보건 복지부
	약사	약사법	약사법 시행령	약사법 시행규칙	대한약사회	보건 복지부
	간호사	간호법	간호법 시행령	간호법 시행규칙	대한 간호협회	보건 복지부
	의료 기사	의료기사 등에 관한 법률	의료기사법 시행령	의료기사법 시행규칙	대한임상 병리사협회 외	보건 복지부

이러한 정의에 따라 본 연구에서 수집하는 뉴스기사는 해당 전문직과 AI를 종합한 키워드(AI + ‘전문직명’)를 한국언론재단 데이터베이스인 빅카인즈(bigkinds)에 검색어로 입력하고, 언론사는 전국일간지⁴⁾로 한정하여, AI가 사회적으로 본격적으로 확대 논의되기 시작한 ‘chatGPT’ 출시일인 2022년 11월 30일부터 연구진행시점인 2026년 1월 31일까지로 기간을 설정하여 검색을 진행하였다. 해당 검색을 진행한 결과 총 9,595개의 기사가 검색되었으며, 각 직종에 따른 신문기사 수 검색결과는 다음 <표 3>와 같다. 이러한 뉴스기사를 통해 워드클라우드, 네트워크 분석, 감성분석을 진행하여 각 전문직종별 분석결과를 제시한다.

<표 3> 전문직별 기사 수

직종	기사 수	직종	기사 수	직종	기사 수
의사	5,840	변리사	100	감정평가사	42
변호사	2,531	한의사	71	법무사	27
간호사	469	치과의사	69	관세사	27
약사	179	노무사	67		
세무사	126	공인회계사	47		

4) 경향신문, 국민일보, 내일신문, 동아일보, 문화일보, 서울신문, 세계일보, 아시아투데이, 조선일보, 중앙일보, 한겨레, 한국일보

다만 상기한 표의 기사 수집 간 ‘AI+전문직명’ 단순 결합은 영역 특화 동의어(리걸테크·법률 AI / 의료 AI·디지털 헬스 / 세무 자동화 등)를 포괄하지 못해 노출 격차가 키워드 범위 차이로 오인될 수 있으나, 해당 부분은 본 연구의 한계로, 다양한 키워드를 반영함으로써 본 연구의 취지보다 범위가 확장되는 한계가 있으므로, 이는 한계로 명시하고, 소비자 인식조사·소셜미디어 텍스트를 보완자료 등은 후속연구에서 반영하기로 한다.

다음으로, AI의 전문직역 기술대체 가능성을 평가하기 위하여 본 연구에서는 각 전문직종에 대한 세부업무를 각 법령 및 O*NET(직무/Task 데이터)를 활용하여 선별한 후 LLM(GPT 5.2, Claude Opus 4.5, Gemini 3 등)을 이용하여 세부업무 별로 대체 가능성을 평가하도록 하였다. 각 직종별 세부업무는 다음 <표 4>⁵⁾와 같다.

<표 4> 전문직 별 세부업무 평가항목

변호사	변리사	회계사	세무사	법무사
【법률 문서 검토 및 작성】	【선행기술 조사】	【회계 감사】	【세무신고 대리】	【등기·등록 신청】
계약서 검토 및 리스크 분석	검색전략 수립(키워드·IPC)	감사계획·위험평가 문서화	신고자료(증빙) 수집·점검	신청서 작성
판례 검색 및 법리 적용	특허·논문 DB 검색/수집	감사절차 수행·테스트 기록	신고서 작성(부가/소득/법인)	첨부서류 준비
법령 적합성 검토	요약 및 대비표(차별점) 작성	감사조서·증적 정리	전자신고 제출·오류 보완	등기소 제출 및 접수
최종 의견서 작성	선행기술 조사보고서 작성	감사보고서 초안·Q&A 정리	정정신고·사후 문의 대응	진행 상황 확인
【판례 검색 및 분석】	【명세서 작성】	【재무제표 작성 및 검토】	【세액 계산】	【서류 작성】
판례 데이터베이스 검색	발명 구성·도면 설명 정리	조정분개 검토·정리	과세표준 산정·조정	정관 작성
판례 요지 추출	청구항 초안 작성	재무제표 작성	공제·감면 적용 검토	계약서 작성
판례 경향 분석	명세서 본문(실시예 등) 작성	주식 작성·정합성 점검	세액 산출근거표 작성	위임장·동의서 작성
사안 적용 가능성 판단	보정/보완 문구 작성	공시서류 검토·제출 준비	납부세액 검증·확정	공증서류 준비

5) 이외 8개 전문직에 대한 세부업무는 부록에서 제시하도록 한다.

【법률 해석 및 자문】	【특허성 판단】	【세무 상담】	【세무조사 대응】	【법률 문서 검토】
법령 조문 해석	신규성 판단·대비표 작성	세무 쟁점 검토	조사 대응자료 준비	등기부등본 분석
사례 분석 및 의견 제시	진보성 판단·근거 정리	세무조정·신고서 검토	소명서·확인서 작성	계약서 법적 검토
클라이언트 상담	기재요건·명확성 검토	세무리스크 진단	질의응답 정리·기록	하자 및 리스크 파악
전략적 법률 조언	거절이유 대응논리 작성	유권해석/질의서 작성	불복절차 서면 초안	개선 의견 제시
【소송 전략 수립】	【특허 전략 수립】	【경영 자문】	【세무 상담】	【법률 상담】
사실관계 정리 및 쟁점 도출	출원 포트폴리오·일정 수립	원가·수익성 분석	세무기초(장부/증빙) 상담	등기 절차 안내
증거 분석 및 평가	우선권·분할·PCT 전략	예산·사업계획 수립 지원	상속·증여 상담	권리관계 설명
소송 시나리오 설계	특허맵·회피설계(FTO) 분석	재무모델·가치평가(DCF)	거래유형 세무처리 안내	분쟁 예방 조언
승소 가능성 평가	라이선스/기술이전 자료 작성	KPI·경영보고서 작성	사전답변/질의서 작성	사후 관리 안내
【법정 변론 및 협상】	【심판·소송 대리】	【내부통제 평가】	【절세 전략 수립】	【소송 서류 작성】
변론 준비서면 작성	불복심판 청구서·이유서 작성	프로세스 문서화(Flow/RCM)	법인전환·급여/배당 설계	소장·답변서 작성
법정 구두변론	무효/정정심판 서면 작성	통제설계 평가	비용처리·감가상각 최적화	증거자료 정리
상대방과의 협상	증거자료 정리·제출	운영테스트·증적 수집	가업승계·증여 전략 검토	준비서면 작성
즉석 법률 대응	구술심리/준비서면 준비	개선권고·평가보고서 작성	연간 절세 실행계획 수립	소송 진행 보조

상기한 표와 같이 본 연구에서는 전문직종의 AI 대체가능성을 직종 단위가 아니라 세부 업무(task) 단위에서 평가한다. 이러한 접근은 기술 변화의 영향이 직업 전체에 균등하게 작용하기보다 직무를 구성하는 과업의 속성(규칙성·반복성·자료구조화 가능성 등)에 따라 이질적으로 나타날 수 있기 때문이다. 따라서 본 연구는 직종별 대표 업무를 세부업무로 분해하고, 각 작업이 AI에 의해 수행(또는 보조)될 수 있는 정도를 코딩(coding)하는 방식으로 대체가능성을 추정한다. 세부업무 평가는 대규모 언어모델(LLM) 3종(GPT 5.2,

Claude Opus 4.5, Gemini 3 Pro)을 자동 코더(annotator)로 활용하여 수행한다.

최근 사회과학 및 커뮤니케이션 연구에서는 LLM이 텍스트 분류·프레임 분류·행위자 식별 등 다양한 코딩 과업에서 인간 코더(특히 크라우드 코더)와 비교 가능한 성능 및 일치도를 보일 수 있음이 보고되었으며, 특히 ChatGPT 계열 모델이 텍스트 주석(annotation) 과업에서 크라우드 워커를 상회하는 정확도를 보였다는 결과는 LLM이 사람이 하던 코딩 업무의 일부를 대체·보조하는 연구 설계가 경험적으로 가능함을 시사한다(Gilardi et al., 2023).

다만, 본 연구에서 LLM이 수행하는 역할은 단순 분류를 넘어, 세부 작업에 대해 평가척도(표준화·확률화 또는 대체가능성 점수)를 부여하는 판정 성격을 포함한다는 점에서 추가적 절차가 요구된다(Gu et al., 2024). 이에 본 연구는 LLM-as-a-Judge 연구에서 권고되는 절차(명시적 평가표 제시, 기준 예시 제공, 다회 반복 평가 및 불일치 점검)를 적용해 평가의 일관성과 설명가능성을 확보하였으며(Gu et al., 2024; Szymanski et al., 2025), 프롬프트·모델·상황에 따른 판정 변동 가능성을 고려해 단일 응답에 의존하지 않고 3종 모델로부터 다중 샘플링한 뒤 합의 규칙(평균·다수결)으로 측정오차를 축소하였다. 프롬프트는 각 세부작업에 대해 (1) 작업 설명, (2) 표준화·확률화 정의, (3) 점수 부여 기준과 판단 근거를 포함하도록 구성하였으며, 6개 평가차원(표 5)에 점수를 부여한다.

〈표 5〉 세부항목별 LLM 기반 평가 예시

예시) 변호사 - [법률 문서 검토 및 작성](전체 업무 비중25%)										
세부작업	비중 (%)	규칙 /20	반복 /20	데이터 /20	기술 /20	예외 /10	검증 /10	합계 /100	AI 대체%	기여6%)
계약서 검토 및 리스크 분석	40%	16 ⁷⁾	16	18	17	5	8	80	80%	32.0%
판례 검색 및 법리 적용	30%	16	16	18	17	5	8	80	80%	24.0%
법령 적합성 검토	20%	20	16	18	17	6	8	85	85%	17.0%
최종 의견서 작성	10%	16	16	18	11	5	8	74	74%	7.4%

차원	배점	핵심 질문	이론 근거
규칙	20	결정규칙이 체크리스트/결정트리로 명시 가능한가	Autor et al.(2003)
반복	20	유사 케이스가 고빈도로 반복되는가	김가봉(2019)
데이터	20	학습·검색·자동처리용 구조화 데이터가 충분한가	Brynjolfsson et al.(2018)
기술	20	현 상용 AI로 현업 수준 구현이 가능한가	Eloundou et al.(2024)
예외	10	예외·특수상황이 적은가(적을수록 고점)	Withey et al.(1983)
검증	10	결과의 객관적 검증·감사가 가능한가	Kroll(2015)

* 배점은 표준화·확률화의 1차 구성차원(규칙·반복·데이터·기술 각 20점)과 보정차원(예외·검증 각 10점)을 구분하여 부여하였다. 표준화 산식은 규칙성을, 확률화 산식은 데이터 학습가능성을 각 개념의 핵심 차원으로 보아 최대 가중치(35)를 두고 예외·검증은 보정차원으로 낮은 가중치를 부여하였다. 가중치 선택의 영향력을 점검한 등배점 대비 민감도 분석에서 직종 순위 상관성이 매우 높아($\rho \approx 0.99$) 결과는 가중 방식에 강건하였다.

이때, 상세 평가에 대한 평가항목을 살펴보면, 규칙의 경우, 규칙을 따르며 수행되는 과업은 컴퓨터가 대체/보완하기 쉽다는 고전적 과업기반 논의와 직접 연결되며, 해당 점수는 법·규정·서식·계산식·체크리스트로 절차가 기술될 수 있는 정도라는 표준화의 핵심 차원으로 해석된다(Autor et al., 2003). 마찬가지로 반복의 경우에도, 과업이 유사 케이스로 반복될수록 규칙화와 자동화(특히 RPA/워크플로우 자동화)의 기대효과가 커지고, 반대로 케이스마다 새롭고 변동이 큰 과업은 대체가 어렵다는 것이 과업기반 변화 연구와 연계된다(김가봉, 2019). 데이터의 경우에는 AI(머신러닝/LLM 포함)이 과업을 학습하려면 데이터가 존재하고, 목적(정답/성과)을 정의할 수 있어야 하며, 오류/성과가 측정되어 개선될 수 있어야 한다는 학습가능성에 기반하여 데이터에 기반한 경우에 AI 대체 가능성이 높아질 수 있음을 확인할 수 있다(Brynjolfsson et al., 2018).

기술의 경우 현재의 구현가능성을 평가하며, LLM이 어떤 과업에 어느 정도 영향을 줄 수 있는지 평가표 기반으로 직업·과업을 분류한 연구는, 평가가 단지 추측이 아니라 기술 능력·과업요구의 대응관계를 정량화할 수 있음을 확인할 수 있다(Eloundou et al., 2024). 또한 예외의 경우, 과업 불확실성을 예외(exceptions)와 연결해 측정하는 조직 기술 연구와 정합성을 가지고 있으며, 비정형·비루틴 영역이 자동화에 저항한다는 논의는 예외 차원을 측정하는 것이 중요하다고 확인할 수 있다(Autor, 2015; Withey et

6) 비중과 AI 대체가능성의 곱

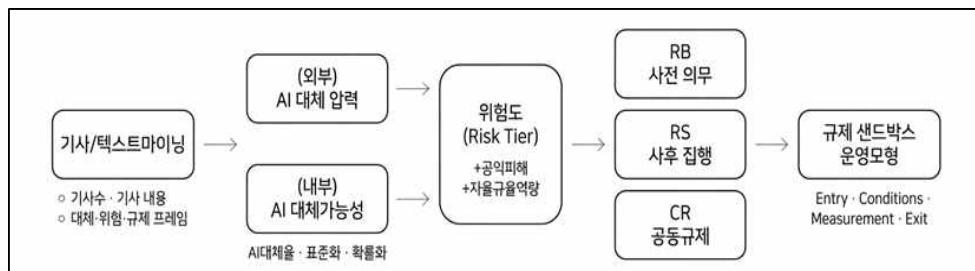
7) 해당 표 내용의 각 점수를 LLM을 통해 도출한다.

al., 1983). 마지막으로 검증의 경우, 정당성 및 감사가능성을 평가하는 지표로 AI가 개입하는 과업은 결과의 정당성·책임성을 위해 검증·감사 가능성이 필수이고, 이는 알고리즘 책임성/감사가능성 논의의 핵심으로 포함되어야 하므로 해당 평가를 실시하였다(Kroll, 2015).

2. 연구모형

본 연구는 인공지능 시대 전문직종의 대체가능성을 단순한 기술 성능의 문제가 아니라, (1) 직무가 AI가 다루기 쉬운 형태로 변환될 수 있는 구조적 조건과 (2) 그 과정에서 형성되는 사회적 압력이 결합된 결과로 파악한다. 이에 따라 연구모형은 (가) 뉴스 데이터에 기반한 외부 압력 측정, (나) 직무 구조에 기반한 내부 대체가능성 평가, (다) 공익피해 및 자율규율 역량을 고려한 위험도(Risk Tier) 산출, (라) 위험도에 비례하는 규제수단 조합(사전 의무 - 사후 집행 - 공동규제), (마) 이를 구현하는 규제샌드박스 운영모형(Entry-Conditions-Measurement-Exit)으로 구성된다. 이를 설명하는 연구모형은 다음 <그림 1>과 같다.

<그림 1> 연구모형



첫째, 외부 AI 대체 압력(사회적 압력)은 기술 도입 과정에서 규제정책 결정을 좌우하는 사회적 수용성과 정치적 파급을 반영한다. 본 연구는 BigKinds 중앙일간지 뉴스 자료를 활용하여 AI와 전문직 대체 이슈에 대한 기사량과 논조를 계량화하고, 필요 시 텍스트마이닝을 통해 프레임(예: 일자리 위협, 혁신·효율, 안전·책임 등)의 분포를 도출한다.

이 변수는 기술적 위험이 실제로 얼마나 큰가와 별개로, 사회가 해당 이슈를 얼마나 민감하게 인지하고 규제 개입을 요구하는가를 나타내는 외생적 압력으로 해석된다.

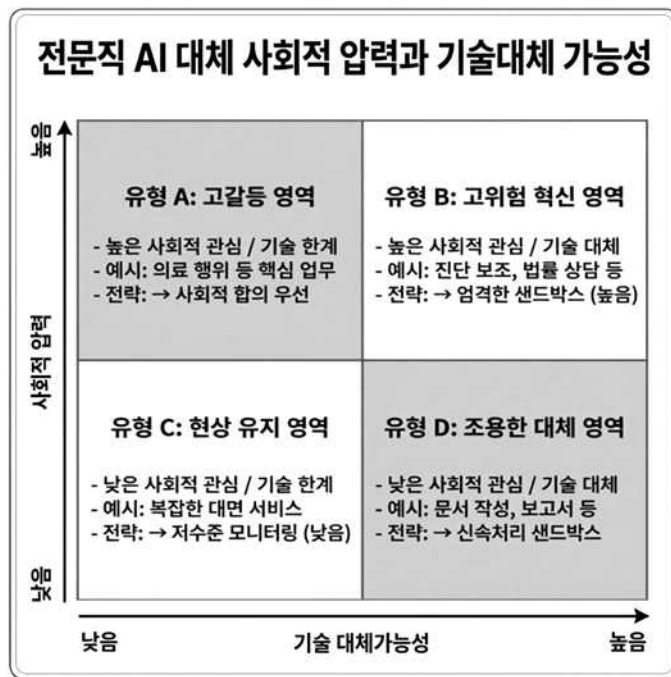
사회적 압력의 분류 임계치는 노출(기사량)과 갈등(부정 논조 %)을 각각 z-표준화한 뒤 노출:갈등 = 3:1 가중으로 종합점수를 산출하고, 표본 내 3분위(tertile)를 사전 임계치로 하여 높음/중간/낮음으로 분류한다. 노출에 우위 가중을 둔 것은 기사량을 의제 강도의 1차 지표로 규정했기 때문이다. 기술축의 고/저(高/低) 이분(二分) 임계치(평균 60%)는 AI 대체율 분포의 최대 자연절단(감정평가사 65.5%-변리사 58.8% 사이 6.7%p 간극)에 근거한다. 또한 RB:RS:CR 권장비율은 추정 모수가 아니라 '위험유형→규제수단' 매핑 설계원리에 따른 예시값으로, 사회적 수용성 위험이 큰 A는 CR을, 안전·책임 위험이 큰 B는 RB를, 확산은 빠르나 주목이 낮은 D는 RS의 비중을 높이는 규칙에서 도출한다.

둘째, 내부 AI 대체가능성(기술적 포획 가능성)은 직무가 AI에 의해 수행·대체될 수 있는 구조적 조건을 측정한다. 본 연구는 이를 표준화와 확률화 두 축으로 구성한다. 표준화는 업무가 규칙·절차·서식·체크리스트 형태로 코드화되어 결과가 재현 가능한 정도를 의미하며, 확률화는 업무가 데이터로 관측되고 학습 목표(정답/성과)와 피드백이 정의되어 AI가 성능을 평가·개선할 수 있는 학습가능성의 정도를 의미한다. 즉 표준화는 업무를 규칙과 절차로 번역하는 과정이고, 확률화는 업무를 데이터와 학습문제로 번역하는 과정이며, 두 조건이 높을수록 AI는 보조적 기능을 넘어 전문직 업무의 핵심 기능을 점진적으로 포획할 가능성이 커진다. 내부 대체가능성은 직종별 업무를 세부업무 단위로 분해하여 LLM을 이용하여 표준화·확률화 점수(및 종합지수)를 산출함으로써 추정한다.

셋째, 본 연구는 외부 압력과 내부 대체가능성을 단순 병렬로 제시하는 데 그치지 않고, 이를 규제결정에 직접 투입 가능한 단일 기준으로 전환하기 위해 위험도를 산출한다. 위험도는 ①기술적 대체가능성(위험의 가능성), ②사회적 압력(위험의 사회적 파급력)을 기본 축으로 하되, 규제의 목적이 공익 보호와 혁신 촉진 간 균형이라는 점을 고려하여 ③공익피해(오류 발생 시 생명·안전·권리 침해 및 사회적 파장)와 ④자율규율 역량(전문직 내부의 윤리·표준·교육·징계 및 준수관리 능력)을 보정요인으로 포함한다. 결과적으로 위험도는 기술이 대체할 수 있느냐와 사회가 민감하게 반응하느냐를 동시에 반영하면서, 고위험·고갈등 영역에서 필요한 규제 강도와 운영 방식의 차등화를 가능하게 한다.

정리하면, 상기한 바와 같이 분석을 진행한 이후, 다음 <그림 2>와 같이 전문직에 대한 AI대체에 대하여 사회적 압력과 기술대체 가능성을 사회적압력(고-저), 기술 대체 가능성(고-저) 2X2 사분면에 위치시킴으로써 유형화하여, 이를 규제샌드박스와 연결시키는 <그림 2>의 개념적 구성으로 사용한다.

<그림 2> 전문직 AI 대체 사회적 압력과 기술대체 가능성에 따른 유형화



넷째, 위험도에 기초하여 규제 개입은 사전 의무(Rule-Based Regulation, 이하 RB), 사후 집행(Reactive Supervision, 이하 RS), 공동규제(Co-Regulation, 이하 CR)의 세 수단을 조합하는 방식으로 설계된다. 사전 의무는 시장 진입 이전 단계에서 최소 안전기준을 확보하기 위한 규제로서, 등록·문서화·검증·설명가능성 확보·인간감독·로그 보존·책임주체 명확화 등과 같은 요건을 포함한다. 사후 집행은 실제 운영 과정에서 위험을 관찰하고 통제하는 규제로서, 사고·민원 보고, 정기 모니터링, 표본감사, 위반 시 시정명령·제재, 리콜·중단 등을 포함한다. 공동규제는 정부와 전문직 단체, 기업, 이용자

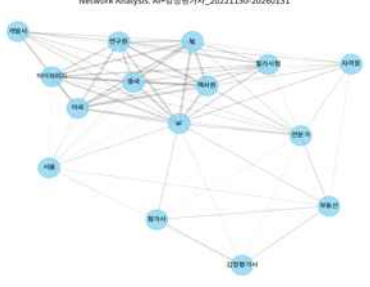
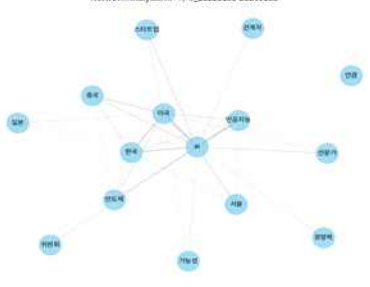
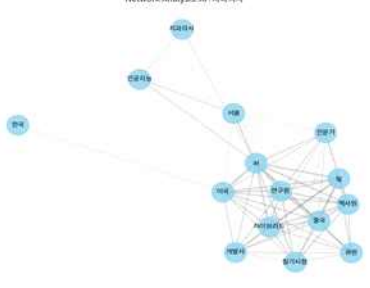
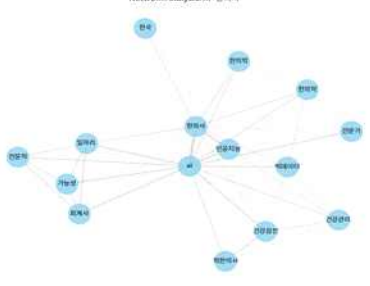
대표 등이 규칙·표준·인증·교육·자율점검 체계를 공동으로 설계·집행하는 방식으로, 전문직의 현장 적합성과 수용성을 확보하는 동시에 포획 위험을 줄이기 위한 투명한 거버넌스 장치(공개, 이해상충 관리, 성과평가)를 전제로 한다. 세 수단은 배타적 선택지가 아니라 위험도에 따라 비중이 달라지는 규제 포트폴리오로 기능하며, 규제샌드박스로 연계한다.

마지막으로, 본 연구의 규제샌드박스 운영모형은 사전 의무(RB), 사후 집행(RS), 공동 규제(CR)의 조합을 ‘진입: Entry-운영조건: Conditions-측정 및 평가: Measurement-종료 및 전환: Exit’의 4단계 운영 규칙으로 구체화한다. 진입 단계에서는 위험도가 높을수록 사전 심사·등록·검증 등 사전 의무(RB) 요건을 강화하여 진입 기준을 상향한다. 운영조건 단계에서는 적용 범위(업무 단계·대상·환경)를 제한하고 인간감독, 고지·동의, 데이터 관리, 책임보험 등의 조건을 부과하여 실증을 통제된 환경에서 수행하도록 한다. 측정 및 평가 단계에서는 오류율·안전사고·편향·민원·피해구제 성과·전문가 검토 이력 등 핵심 성과지표를 정해 정기 보고와 검증을 수행하고, 사후 집행(RS)을 통해 운영 중 위험을 교정한다.

종료 및 전환 단계에서는 사전에 설정한 임계치와 중단 기준에 따라 ① 제도권 편입(확대), ② 조건 강화 후 연장, ③ 실증 중단·퇴출을 결정함으로써 규제의 학습과 조정을 제도화한다. 세 규제수단의 위험유형별 권장비용은 통계적 추정값이 아니라, 위험유형의 성격과 각 규제수단의 기능을 매핑한 설계 규칙(design rule)에 따른 예시값이다. 구체적으로 본 연구는 (1) RB(사전 의무)는 안전·책임 위험이 클수록, (2) RS(사후 집행)는 확산 속도가 빠르거나 운영 중 위험 관찰 필요가 클수록, (3) CR(공동규제)은 사회적 수용성·정당성 위험이 클수록 비중을 높인다는 세 가지 매핑 원칙을 설정한다. 이러한 운영 설계는 불확실성이 큰 기술환경에서 일률적 규제 대신 위험과 증거에 비례하여 규제 강도를 조정하는 유연한 거버넌스를 구현한다는 점에서, 본 연구가 제시하는 위험도 기반 규제 설계의 실행 메커니즘이 될 수 있다.

세무사		
법무사		
공인노무사		
관세사		

8) 긍정-부정에 대한 판정 후 중립 논조가 50%가 넘는 경우, 중립을 병기함

감정평가사		
의사		
치과의사		
한의사		

약사																																																										
간호사																																																										
의료기사	기사 수가 매우 적으므로 분석하지 않음	기사 수가 매우 적으므로 분석하지 않음																																																								
직종별 뉴스기사 감성분석	직종별 AI 뉴스 감성 분석 (13개 직종) <table><thead><tr><th>직종</th><th>공정 (%)</th><th>부정 (%)</th><th>중립 (%)</th></tr></thead><tbody><tr><td>법무사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>변호사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>감정평가사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>한의원</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>의사</td><td>29.10</td><td>12.50</td><td>58.40</td></tr><tr><td>공인회계사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>간호사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>약사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>세무사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>노무사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>치과의사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>관세사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr><tr><td>변리사</td><td>23.20</td><td>16.00</td><td>60.80</td></tr></tbody></table>		직종	공정 (%)	부정 (%)	중립 (%)	법무사	23.20	16.00	60.80	변호사	23.20	16.00	60.80	감정평가사	23.20	16.00	60.80	한의원	23.20	16.00	60.80	의사	29.10	12.50	58.40	공인회계사	23.20	16.00	60.80	간호사	23.20	16.00	60.80	약사	23.20	16.00	60.80	세무사	23.20	16.00	60.80	노무사	23.20	16.00	60.80	치과의사	23.20	16.00	60.80	관세사	23.20	16.00	60.80	변리사	23.20	16.00	60.80
직종	공정 (%)	부정 (%)	중립 (%)																																																							
법무사	23.20	16.00	60.80																																																							
변호사	23.20	16.00	60.80																																																							
감정평가사	23.20	16.00	60.80																																																							
한의원	23.20	16.00	60.80																																																							
의사	29.10	12.50	58.40																																																							
공인회계사	23.20	16.00	60.80																																																							
간호사	23.20	16.00	60.80																																																							
약사	23.20	16.00	60.80																																																							
세무사	23.20	16.00	60.80																																																							
노무사	23.20	16.00	60.80																																																							
치과의사	23.20	16.00	60.80																																																							
관세사	23.20	16.00	60.80																																																							
변리사	23.20	16.00	60.80																																																							
기사 수 및 감성분석(공정-중립-부정)에 따른 뉴스기사 분석결과																																																										
직종	분석 기사 수	공정 (%)	부정 (%)	중립 (%)	종합 톤8)																																																					
의사	5,840	29.10%	12.50%	58.40%	공정/중립																																																					
변호사	2,531	23.20%	16.00%	60.80%	공정/중립																																																					

간호사	469	34.80%	10.70%	54.60%	긍정/중립
약사	179	35.20%	11.20%	53.60%	긍정/중립
세무사	126	35.70%	11.10%	53.20%	긍정/중립
변리사	100	46.00%	10.00%	44.00%	긍정
한의사	71	26.80%	16.90%	56.30%	긍정/중립
치과의사	69	42.00%	10.10%	47.80%	긍정
공인노무사	67	35.80%	23.90%	40.30%	긍정
공인회계사	47	31.90%	17.00%	51.10%	긍정/중립
감정평가사	42	26.20%	7.10%	66.70%	긍정/중립
법무사	27	14.80%	22.20%	63.00%	부정/중립
관세사	27	44.40%	14.80%	40.70%	긍정

- 워드클라우드 파라미터: generate_from_frequencies
- 네트워크 분석 파라미터: 키워드 선별 (top_keywords=15), 노드 설정 (node_size=1500, node_color='skyblue', alpha=0.8), 에지(선) 가중치 (width=scaled_weights)(5 * weight / max_weight)
- 긍정: 기사에서 AI 도입을 효율화·혁신·기회로 평가하거나, 해당 직종에 우호적 전망을 제시하는 경우
- 부정: 기사에서 AI에 의한 일자리 위협·대체·전문성 약화·갈등을 부각하거나, 해당 직종에 위협적 전망을 제시하는 경우
- 중립: 사실 전달 위주이거나 긍·부정 단서가 혼재·상쇄되어 어느 방향으로도 판정하기 어려운 경우
- 뉴스 논조(긍·부정·중립)는 상용 대규모 언어모델(LLM)을 분류기(Claude(Anthropic), 버전 Opus 4.7)로 활용하여 기사 단위로 평가
- LLM에는 (1) 분석 과업 설명, (2) 위 3범주의 정의와 anchor 예시, (3) 출력 형식(범주 1개와 판단 근거)을 포함한 프롬프트를 제시(본문+제목, 중립밴드 ±0.7)

뉴스 기사를 분석한 결과, 의사와 변호사 직군이 전체 논의의 약 80% 이상을 차지하며 AI 도입에 대한 사회적 관심이 가장 집중되어 있음으로 도출되었다. 먼저, 워드클라우드 및 네트워크 분석결과 보건 의료 계열의 경우(의사, 간호사, 치과의사, 약사), AI 의료 기기, 진단 보조 솔루션이 이미 주류 시장에 진입했음을 보여주며, AI가 과도한 업무를 줄여 환자 케어에 집중하게 만든다는 업무 효율화 관점이 지배적임을 의미할 수 있다. 법률 및 지식재산의 경우, 서면 작성 등 혁신 기사와 판례 분석 알고리즘에 대한 규제 논의가 공존하며 보조적 공존을 모색하고 있으며, 변리사의 경우에는 AI가 생성한 발명의 특허 인정 여부 등 미래지향적 논의가 많으며, 특히 심사 효율화를 국가 경쟁력과 연결하는 혁신 프레임이 압도적으로 보인다.

회계 및 세무, 노무 계열의 경우, 복잡한 세무 계산을 자동화하는 서비스들이 출시되면

서 전문가의 업무 편의성이 증대되었다는 긍정적 기류가 있는 반면에, AI가 회계 감사를 수행하면서 인력 수요가 줄어든 것이라는 우려가 실제 수험생들의 기피 현상으로 이어지는 양상을 보도하거나, 윤리적 문제와 노무사 고용 시장의 위축이 갈등 프레임을 보여주는 등의 보도가 이루어졌다. 평가 및 통관 등의 계열에서는 통관 절차 자동화로 인한 수출입 경쟁력 강화라는 국가 산업적 측면에서 긍정적인 프레임이 형성하거나 주관적 평가가 개입될 수 있는 영역에 AI를 도입해 투명성을 확보하려는 산업 인프라 혁신 성격이 강하였다.

감성 분석 측면에서는 변리사와 관세사의 경우, AI 관련 뉴스기사의 긍정성이 높았으며, 법무사와 노무사는 상대적으로 부정적인 우려가 있음을 확인할 수 있었다. 특히, 의료(의사/간호사)와 법률(변호사) 분야의 AI 도입 논의가 압도적이었으며, 이는 생명 및 권리와 직결된 분야인 만큼 기술 도입에 따른 사회적 파급력이 크기 때문으로 분석되었다. 또한 변리사, 관세사, 치과의사 등 기술적 도구를 활용하여 정밀도나 효율성을 높일 수 있는 직종에서는 AI를 '혁신적 도구'로 인식하는 경향이 뚜렷하였다. 법무사, 노무사, 회계사 등 상담 및 서류 작성이 주 업무인 직종에서는 AI에 의한 '업무 대체', '수험생 감소', '전문성 약화' 등의 키워드가 부정적 평가를 견인하고 있음을 확인하였다. 기사에 대한 언론 프레임링을 분석해보면 다음 <표 7>과 같다.

<표 7> 뉴스기사의 발행 의도 및 언론 프레임링

직종	기사 발행 의도 (Intent)	언론 프레임링 분석 (Framing Analysis)
의사	의료 데이터 고도화 및 정밀 진단 보조 기술 도입 홍보	(보조적 공존) AI를 의료진의 전문성을 강화하는 보조자로 설정
변호사	리걸테크 발달에 따른 법률 서비스 접근성 및 혁신 보고	(보조적 공존) 법률 전문가와 AI의 협업 모델 및 윤리적 가이드 제시
간호사	돌봄 노동 경감 및 실시간 환자 모니터링 효율화 강조	(업무 효율화) 단순 반복 업무 자동화를 통한 서비스 질 향상
약사	의약품 처방 정확도 제고 및 스마트 약국 시스템 구축 현황	(업무 효율화) AI를 통한 오투약 방지 및 고객 편의 증대
세무사	세무 행정 전산화 및 납세자 편의 기능 강화 성과 공유	(업무 효율화) 수작업 고충 해결 및 정확한 세무 서비스 제공
변리사	지식재산권(IP) 심사 효율화 및 국가 기술 경쟁력 강화	(혁신 및 기회) 기술 선도와 지식재산 강국 도약의 동력으로 묘사

한 의사	한방 진단 보조 장비의 디지털화 및 데이터 표준화 시도	(업무 효율화) 한의학의 객관적 데이터 확보를 통한 신뢰도 제고
치과 의사	디지털 덴티스트리 도입을 통한 진료 정밀도 확보	(혁신 및 기회) 전통적 진료 방식의 디지털 전환 성과 강조
공인 노무사	노사 갈등 중재 및 상담 시 AI 활용의 실효성 논의	(위협 및 갈등) AI 도입에 따른 고용 불안 및 시장 혼란 강조
공인 회계사	AI 도입에 따른 회계 감사 환경 변화 및 수험 시장 영향	(불안 및 생존) 고용 시장 위축과 직종 매력도 하락에 집중
감정 평가사	부동산 가치 평가 모델의 자동화 및 데이터 신뢰성 확보	(업무 효율화) 평가 프로세스의 투명화와 데이터 중심 전환
법무사	등기/서류 업무 자동화에 따른 직무 영역 변화 보고	(위협 및 갈등) 전문직 업무 영역 침해와 대체 가능성에 대한 우려
관세사	스마트 물류 체계 구축 및 관세 행정 투명성 제고	(혁신 및 기회) 물류 산업 인프라 혁신의 핵심 기술로 정의

상기한 바와 같이 기사 수가 가장 많은 의사와 변호사 직군은 AI를 경쟁자가 아닌 전문성을 보완하는 동반자로 프레이밍하는 경향이 강하고, 이는 기술적 완성도와 별개로 책임 소재 등 인간 전문가의 최종 판단이 중요한 분야이기 때문으로 보인다. 또한 혁신과 성장(변리사/관세사/치과)의 경우, 기술 도입의 성과가 뚜렷한 직종으로 AI를 산업 경쟁력을 높이는 기회로 프레이밍하며 긍정적인 논조로 보인다. 위협과 대체(노무/법무/회계)의 경우, 업무의 정형화 정도가 높거나 상담 비중이 큰 직종으로 AI 도입을 전문직의 위기 또는 시장 잠식 프레임으로 다루는 경향이 있어, 수험생 감소나 고용 불안 이슈와 긴밀하게 연결되어 있다.

결과를 종합해 보면, 전문성(Judgement)이 강조되는 직종(의사, 변호사, 변리사)은 AI를 파트너로 인식하여 긍정적인 담론을 형성하는 반면, 정형화된 업무(Process) 비중이 큰 직종(법무사, 노무사, 회계사)은 AI를 생존의 위협으로 인식하는 경향이 뉴스 기사를 통해 확인할 수 있다.

2. 전문직에 대한 AI 대체가능성 분석결과

먼저, 3종 모델의 직종별 산출값에 대해 〈표 8〉과 같이 급내상관(ICC)·Krippendorff α ·Kendall W를 산출하였다. 모델 간 체계적 수준편향(Gemini가 타 모델 대비 평균 약 11%p 높게 평가)을 통제하기 위해 각 모델 점수를 z-표준화한 값을 신뢰도 보고 기준으로 삼는다. 그 결과 AI 대체율 $ICC(C,k)=.95 \cdot ICC(A,k)=.95 \cdot \alpha=.87$, 표준화 $ICC(C,k)=.96 \cdot \alpha=.90$, 확률화 $ICC(C,k)=.76 \cdot \alpha=.53$ 로, 세 지표 모두 일관성 기준 양호~우수 수준의 신뢰도를 확보하였다(Koo & Li, 2016). 확률화는 세 지표 중 신뢰도가 가장 낮으나 z-표준화 후 $ICC(C,k)=.76$ 의 양호한 일관성을 확보하므로, 표준화와 함께 직무구조를 구성하는 두 축으로 활용한다.

〈표 8〉 LLM 판정의 측정 신뢰도 결과

측정치	ICC(C,k)	ICC(A,k)	Krippendorff α	Kendall W	판정
AI 대체율	.95	.95	.87	.93	우수
표준화	.96	.96	.90	.95	우수
확률화	.76	.77	.53	.66	양호

또한 LLM 판정의 측정 신뢰도를 모델 간 일치도(ICC·Krippendorff α ·Kendall W)로 검증한 데 더하여, 판정 결과의 외부 수렴타당도(convergent validity)를 선행 전문가 연구와의 대조를 통해 확인하였다. 전 직종을 포괄하는 인간 전문가 코더의 신규 섭외는 14개 전문직종 전체에 걸친 동일 평가표 기반 평가가 가능한 패널 구성의 현실적 제약 때문에 본 연구의 범위에서 수행하기 어려운 측면이 있다. 이에 본 연구는 장주희 외(2020)의 의사 직무에 대한 FGI·심층면담·델파이 기반 평가, 송성수(2022)의 직종별 대체확률 추정, 김건우(2018)가 Frey & Osborne(2017; 2013 워킹페이퍼) 방법론을 한국 423개 직업에 적용해 산출한 대체확률 등 이미 인간 전문가의 평가 절차를 거쳐 공표된 선행연구의 결과값을 외부 비교 기준으로 활용하였다. 이러한 접근은 신규 전문가 패널 구성 시 불가피한 표본 한계(직역 편향·소수 의견)에 비해 다층적 전문가 집단(FGI·델파이·통계 추정)이 공식 절차를 통해 도출한 평가를 비교 기준으로 삼는다는 점에서 외

부 타당도 검증의 보완적 근거로 활용될 수 있다. 해당 구체적 대응성은 <표 13>에 제시한다.

전문직의 AI 대체가능성은 모델 성능 자체보다, 업무가 ① 규칙·절차로 형식화될 수 있는지(표준화), ② 데이터·피드백이 축적되어 학습 문제로 환원될 수 있는지(확률화)의 해 좌우될 수 있다. 이런 관점은 LLM의 노동시장 영향(작업 노출)을 세부업무 수준의 기준으로 추정한 연구들과 일맥을 같이한다(Eloundou et al., 2024). 따라서, LLM(chatGPT 5.2, Claude Opus 4.5, Gemini 3 pro)에게 평가표를 프롬프트로 제시한 후, 해당 표에 근거하여 세부업무 별로 평가항목에 점수를 매겨 AI의 전문직 대체가능성을 평가하였다.

본 연구는 각 세부작업을 6개 항목으로 채점하고(총 100점), 이를 그대로 기술적 대체가능성(AI 대체%)으로 해석한다. 작업 목록 구성 시에는 법령·지침 기반 세부업무(task) 정리와 더불어, 해외 비교/형식 표준을 위해 O*NET처럼 occupation-specific task statement를 제공하는 데이터베이스의 체계를 참조할 수 있다. 직종은 여러 업무영역과 그 안의 세부업무로 구성되므로, 단순 평균이 아니라 업무비중 가중평균으로 보고하며, 이는 하단 <표 9>와 같다.

<표 9> 대체가능성 분석 평가표

평가 항목	배점	세부기준(질문)	점수 범위	증거	주의/해석
규칙	20	결정규칙이 체크리스트/결정트리로 명시 가능한가?	0=직관/경험이 핵심 1=원칙만 존재 2=핵심 일부 문서화 3=대부분 기준이 항목화 4=거의 전부 규칙화	법령/시행령, 내부 SOP, 체크리스트, 심사지침	규칙(/20)=아래 5문항(각 0~4) 합계
		입력값이 정형(필드/코드/수치)으로 구조화되어 있는가?	0=대화·관찰 중심 1=비정형 대부분 2=정형+비정형 혼재 3=정형이 대부분 4=코드/수치/서식 고정	전자신고 필드, 표준서식, 코드체계	
		처리 절차가 표준(SOP/워크플로우)으로 고정되어	0=케이스마다 절차 상이 1=느슨한 절차 2=표준있으나 예외 다수	업무 프로세스 문서, 결재/검토 단계	

		있는가?	3=표준절차 지배적 4=단계·책임·산출물고정		
		산출물이 템플릿/서식 기반인가?	0=산출물 매번 다름 1=형식 느슨 2=템플릿있으나 자유도 큼 3=템플릿기반이일반 4=서식·문구·구성이표준화	보고서/신청서 템플릿, 전자서식	
		근거·판단 매핑이 가능한가(규정/판례/ 가이드로 설명 가능)?	0=근거 제시 곤란 1=일부만 가능 2=절반정도 가능 3=대부분 가능 4=거의 전부 근거와 매핑	조문·판례·심사 지침 인용 관행, 근거표	경계사례(해석·재 량)가 크면 1~2점 감점
반복	20	유사 케이스가 고빈도로 반복되는가(유형 분산 대비 반복성)?	0=거의 매번 다름 1=반복 낮음 2=중간(반반) 3=반복 높음 4=대량·루틴 반복	월 처리건수, 상위 5유형 비중, 동일 서식 반복률	(0~4점)×5 → 0~20 범위로 환산
데이터	20	AI 투입 가능한 데이터가 존재·접근·품질관리 되는가?	0=비디지털/접근 불가 1=부분적 디지털 2=혼재·품질 편차 3=대부분 디지털·검색 가능 4=DB화·라벨/메타 관리	전자문서, DB, 표준코드, 메타데이터	(0~4점)×5 → 0~20 범위로 환산
기술	20	현 시점 상용 기술로 현업 수준 구현이 가능한가(보조→운영) ?	0=현업 구현 난해 1=부분 보조 2=사람 검토 전제 자동화 3=사람 승인형 운영 가능 4=대량·안정 운영가능	필요 기능(OCR/추출 /검색/규칙엔진/ 검증) 상용성	(0~4점)×5 → 0~20 범위로 환산
예외	10	예외·특수상황이 얼마나 많은가(예외가 적을수록 고점)?	0=예외가 본질 2=예외 매우 잦음 4=예외 잦음(범주화일부) 6=예외 있으나 범주화 가능 8=예외 제한적 10=예외 거의없음	재작업률, 상급자 개입 비율, 예외 유형 수	0~10점 범위에서 ±1 조정 허용. '예외 많음'이면 감점(보수적)
검증	10	결과를 객관적으로 맞다/틀리다로 판정·교차검증하기 쉬운가?	0=객관 판정 곤란 2=검증 비용 매우큼 4=부분 검증 6=기준 있으나 해석여지 8=대부분 검증 가능 10=자동 검증 가능 수준	체크리스트, 규정충족, 계산검산, DB 매칭, 이중검수	0~10점 범위에서 ±1 조정 허용. '설득/전략' 중심이면 저점

LLM에게 제시한 평가 6개 항목

- 규칙(20): 규정·서식·산식·체크리스트로 절차를 대부분 표현 가능한가
- 반복(20): 유사 케이스가 얼마나 자주 반복되는가(루틴성)
- 데이터(20): 학습/검색/자동처리를 위한 디지털·구조화 데이터가 충분한가
- 기술(20): 현 상용 AI(LLM/OCR/RPA/CV 등)로 구현 가능한가(정확도·운영)
- 예외(10): 예외·특수상황이 얼마나 많은가(많을수록 감점)
- 검증(10): 결과의 객관적 검증·감사(정답/산식/교차검증)가 가능한가
- 전체비중 = 섹션비중 × 섹션 내 task 비중
- 직종 AI 대체율 = $\sum(\text{전체비중} \times \text{AI}(\text{task}))$
- 표준화_task(0~10, 소수점 반올림)

$$= 10 \times (0.35 \cdot \text{규칙}/20 + 0.20 \cdot \text{반복}/20 + 0.20 \cdot \text{데이터}/20 + 0.15 \cdot \text{예외}/10 + 0.10 \cdot \text{검증}/10)$$
- 확률화_task(0~10, 소수점 반올림)

$$= 10 \times (0.35 \cdot \text{데이터}/20 + 0.25 \cdot \text{반복}/20 + 0.20 \cdot \text{검증}/10 + 0.10 \cdot (1 - \text{규칙}/20) + 0.10 \cdot (1 - \text{예외}/10))$$

상기한 바와 같이 따라서 제시한 평가표에 따라 LLM은 전문직에 대한 AI 대체가능성 세부 평가를 진행한 뒤, 각 평가항목별로 분석결과를 제시하며, AI가 평가표에 따라 평가한 예시는 하단의 <표 10>과 같다.

<표 10> 대체 가능성 평가 상세 점수표 예시(변호사)

변호사 - AI 대체 가능성 상세 분석(GPT)										
AI 대체율:	58.4%	표준화:		5/10		확률화:		5/10		
【법률 문서 검토 및 작성】(비중25%, AI대체80%)										
세부작업	비중%	규칙 /20	반복 /20	데이터 /20	기술 /20	예외 /10	검증 /10	합계 /100	AI%	기여%
계약서 검토 및 리스크 분석	40%	16	16	18	17	5	8	80	80%	32.0%
판례 검색 및 법리 적용	30%	16	16	18	17	5	8	80	80%	24.0%
법령 적합성 검토	20%	20	16	18	17	6	8	85	85%	17.0%
최종 의견서 작성	10%	16	16	18	11	5	8	74	74%	7.4%
합계									79.75	19.9%
【판례 검색 및 분석】(비중20%, AI대체64%)										
세부작업	비중%	규칙 /20	반복 /20	데이터 /20	기술 /20	예외 /10	검증 /10	합계 /100	AI%	기여%
판례	40%	10	10	18	17	5	5	65	65%	26.0%

데이터베이스 검색										
판례 요지 추출	30%	10	10	18	17	5	5	65	65%	19.5%
판례 경향 분석	20%	10	10	18	17	5	5	65	65%	13.0%
사안 적용 가능성 판단	10%	2	10	18	17	5	5	57	57%	5.7%
합계									63	15.8%
【법률 해석 및 자문】 (비중20%, AI대체43%)										
세부작업	비중%	규칙 /20	반복 /20	데이터 /20	기술 /20	예외 /10	검증 /10	합계 /100	AI%	기여%
법령 조문 해석	40%	20	4	18	11	3	3	59	59%	23.6%
사례 분석 및 의견 제시	30%	10	4	10	11	2	3	40	40%	12.0%
클라이언트 상담	20%	2	4	10	4	2	3	25	25%	5.0%
전략적 법률 조언	10%	2	4	10	4	2	3	25	25%	2.5%
합계									37.25	9.3%
【소송 전략 수립】 (비중20%, AI대체25%)										
세부작업	비중%	규칙 /20	반복 /20	데이터 /20	기술 /20	예외 /10	검증 /10	합계 /100	AI%	기여%
사실관계 정리 및 쟁점 도출	40%	2	4	4	4	2	3	19	19%	7.6%
증거 분석 및 평가	30%	2	4	10	11	2	3	32	32%	9.6%
소송 시나리오 설계	20%	2	4	10	4	2	3	25	25%	5.0%
승소 가능성 평가	10%	2	4	10	4	2	3	25	25%	2.5%
합계									25.25	6.3%
【법정 변론 및 협상】 (비중15%, AI대체30%)										
세부작업	비중%	규칙 /20	반복 /20	데이터 /20	기술 /20	예외 /10	검증 /10	합계 /100	AI%	기여%
변론 준비서면 작성	40%	8	10	4	11	2	3	38	38%	15.2%
법정 구두변론	30%	2	10	4	4	2	3	25	25%	7.5%
상대방과의 협상	20%	2	10	4	4	2	3	25	25%	5.0%
즉석 법률 대응	10%	2	10	4	4	2	3	25	25%	2.5%
합계									28.25	7.1%

상기한 바와 같이 분석결과에 따라, 각 AI별 분석결과는 아래 <표 11>과 같다. 표준화, 확률화, AI 대체율은 평가방법에 따르며, 위험도는 60% 이상 높음, 50% 이상 ~ 60% 미만 중간, 50% 미만 낮음으로 평가한다.

<표 11> LLM(chatGPT, Claude, Gemini)-평가 세부항목 기반 평가결과

chatGPT 5.2 thinking 평가결과					
순위	전문직	표준화	확률화	AI대체율	위험도
1	법무사	7/10	6/10	81.40%	높음
2	관세사	7/10	6/10	72.20%	높음
3	의료기사	6/10	7/10	71.80%	높음
4	약사	6/10	6/10	67.80%	높음
5	감정평가사	6/10	6/10	66.70%	높음
6	세무사	6/10	6/10	64.40%	높음
7	공인노무사	6/10	6/10	64.10%	높음
8	공인회계사	6/10	6/10	64.00%	높음
9	변호사	5/10	5/10	58.40%	중간
10	의사	5/10	6/10	54.50%	중간
11	변리사	5/10	5/10	52.90%	중간
12	간호사	5/10	6/10	52.10%	중간
13	한의사	5/10	5/10	50.10%	중간
14	치과의사	4/10	5/10	46.20%	낮음
Claude Opus 4.5 평가결과					
순위	전문직	표준화	확률화	AI대체율	위험도
1	법무사	8/10	6/10	71.1%	높음
2	관세사	8/10	6/10	68.5%	높음
3	약사	7/10	7/10	62.9%	높음
4	세무사	7/10	6/10	62.7%	높음
5	의료기사	7/10	7/10	62.6%	높음
6	감정평가사	7/10	6/10	59.4%	중간
7	공인노무사	7/10	6/10	59.1%	중간
8	공인회계사	7/10	6/10	59.0%	중간
9	변호사	6/10	6/10	51.0%	중간
10	변리사	6/10	6/10	51.0%	중간
11	의사	6/10	6/10	48.8%	낮음
12	간호사	6/10	6/10	47.4%	낮음

13	한의사	5/10	5/10	45.7%	낮음
14	치과 의사	5/10	5/10	41.8%	낮음
Gemini 3 pro					
순위	전문직	표준화	확률화	AI 대체율	위험도
1	관세사	9/10	9/10	86.5%	높음
2	법무사	9/10	8/10	83.1%	높음
3	세무사	8/10	9/10	82.7%	높음
4	약사	8/10	9/10	82.3%	높음
5	공인회계사	8/10	8/10	80.9%	높음
6	공인노무사	8/10	8/10	77.5%	높음
7	변리사	7/10	8/10	72.4%	높음
8	감정평가사	7/10	8/10	70.3%	높음
9	의료기사	7/10	7/10	69.8%	높음
10	변호사	5/10	7/10	60.4%	높음
11	의사	5/10	7/10	57.5%	중간
12	치과 의사	5/10	6/10	52.5%	중간
13	간호사	4/10	4/10	40.3%	낮음
14	한의사	3/10	4/10	31.2%	낮음

LLM 별 종합 판단 결과는 다음 <표 12>와 같다. 전문직종별 AI 대체가능성(이하 AI 대체율)과 직무구조 지표(표준화·확률화)를 단일 모델의 판단에 의존하지 않고, 서로 다른 상용 대규모 언어모델(LLM) 3종(ChatGPT 5.2 Thinking, Claude Opus 4.5, Gemini 3 Pro)의 평가를 병렬적으로 수집한 뒤 앙상블 방식으로 통합하였다. 구체적으로 각 직종 j 에 대해 모델 $m \in \{GPT(1), Claude(2), Gemini(3)\}$ 의 AI 대체율 산출값을 확보하고, 직종별 최종 AI 대체율은 단순평균으로 정의하였다.

$$AI_j^{ens} = \frac{1}{3} \sum_{m=1}^3 AI_{jm}$$

동시에 모델 간 평가 불일치를 불확실성 지표로 보고, 직종별 표준편차를 산출하여(모델 3종의 산출값 분산) 결과 해석에 활용하였다. 표준화 및 확률화(각 0-10점) 역시 동일한 방식으로 산출하였다. 위험도(높음/중간/낮음)는 각 모델이 제시한 범주값을 취합한 뒤 다수결(majority vote)로 최종 범주를 결정하였다(동률 발생 시 상위 위험 범주를 우선).

〈표 12〉 LLM 평가 세부항목 기반 평가결과 다수결 종합

AI 대체율 종합							
순위(평균)	전문직	AI%_GPT	AI%_Claude	AI%_Gemini	AI%_평균	AI%_표준편차	위험도(합의)
1	법무사	81.4	71.1	83.1	78.5	5.3	높음
2	관세사	72.2	68.5	86.5	75.7	7.8	높음
3	약사	67.8	62.9	82.3	71	8.2	높음
4	세무사	64.4	62.7	82.7	69.9	9.1	높음
5	의료기사	71.8	62.6	69.8	68.1	4	높음
6	공인회계사	64	59	80.9	68	9.4	높음
7	공인노무사	64.1	59.1	77.5	66.9	7.8	높음
8	감정평가사	66.7	59.4	70.3	65.5	4.5	높음
9	변리사	52.9	51	72.4	58.8	9.7	중간
10	변호사	58.4	51	60.4	56.6	4	중간
11	의사	54.5	48.8	57.5	53.6	3.6	중간
12	치과 의사	46.2	41.8	52.5	46.8	4.4	낮음
13	간호사	52.1	47.4	40.3	46.6	4.9	낮음
14	한 의사	50.1	45.7	31.2	42.3	8.1	낮음
표준화							
순위(평균)	전문직	GPT	Claude	Gemini	평균		
1	법무사	7	8	9	8		
2	관세사	7	8	9	8		
3	약사	6	7	8	7		
4	세무사	6	7	8	7		
5	의료기사	6	7	7	6.7		
6	공인회계사	6	7	8	7		
7	공인노무사	6	7	8	7		
8	감정평가사	6	7	7	6.7		
9	변리사	5	6	7	6		
10	변호사	5	6	5	5.3		

11	의사	5	6	5	5.3
12	치과의사	4	5	5	4.7
13	간호사	5	6	4	5
14	한의사	5	5	3	4.3
확률화					
순위(평균)	전문직	GPT	Claude	Gemini	평균
1	법무사	6	6	8	6.7
2	관세사	6	6	9	7
3	약사	6	7	9	7.3
4	세무사	6	6	9	7
5	의료기사	7	7	7	7
6	공인회계사	6	6	8	6.7
7	공인노무사	6	6	8	6.7
8	감정평가사	6	6	8	6.7
9	변리사	5	6	8	6.3
10	변호사	5	6	7	6
11	의사	6	6	7	6.3
12	치과의사	5	5	6	5.3
13	간호사	6	6	4	5.3
14	한의사	5	5	4	4.7

상기한 <표 12>는 14개 전문직종에 대해 3개 LLM이 산출한 AI 대체율을 통합한 결과를 제시한다. 양상블 평균 기준 대체가능 상위권은 법무사(78.5%), 관세사(75.7%), 약사(71.0%), 세무사(69.9%), 의료기사(68.1%) 순으로 나타났으며, 그 다음으로 공인회계사(68.0%), 공인노무사(66.9%), 감정평가사(65.5%)가 뒤를 이었다. 중위권은 변리사(58.8%), 변호사(56.6%), 의사(53.6%)였고, 하위권은 치과의사(46.8%), 간호사(46.6%), 한의사(42.3%)로 나타났다.

위험도 합의(다수결) 기준으로는 8개 직종(법무사, 관세사, 약사, 세무사, 의료기사, 공인회계사, 공인노무사, 감정평가사)이 ‘높음’, 3개 직종(변리사, 변호사, 의사)이 ‘중간’, 3개 직종(치과의사, 간호사, 한의사)이 ‘낮음’으로 분류되었다. 이는 대체가능성이 높게 평가된 직종이 위험도 역시 높게 판정되는 경향을 보이나, 일부 직종에서는 모델 간 판단 차이가 존재함을 시사한다.

표준화는 양상블 평균 기준 법무사·관세사(각 8.0/10)가 가장 높고, 약사·세무사·공인 회계사·공인노무사(각 7.0/10)가 그 다음으로 나타났다. 확률화는 약사(7.3/10)가 가장 높고, 관세사·세무사·의료기사(각 7.0/10)가 뒤를 이었다. 반면 한의사(표준화 4.3/10, 확률화 4.7/10), 치과의사(표준화 4.7/10, 확률화 5.3/10) 등은 상대적으로 낮게 평가되었다. 이는 직종별로 업무의 절차·서식화 수준(표준화)과 데이터·피드백 기반의 학습가능성(확률화)이 상이하며, 해당 구조적 조건이 AI 대체율의 차이를 설명하는 핵심 요인으로 기능할 수 있음을 시사한다.

3. 종합결과 및 규제유형과의 연결

상기한 결과를 종합하여 하단의 <표 13>은 본 연구의 핵심 분석틀을 구성하는 두 축, 즉 기술적 대체가능성(내부 요인)과 사회적 압력(외부 요인)을 직종별로 결합하여 위험도(A-D)를 도출하고, 각 위험유형에 부합하는 규제 포트폴리오(RB-RS-CR) 및 규제샌드박스 운영규칙(Entry-Conditions-Measurement-Exit)을 제시한다. 먼저 기술적 대체가능성은 단일 모델의 주관적 판단에 의존하는 위험을 줄이기 위해, ChatGPT 5.2 Thinking, Claude Opus 4.5, Gemini 3 Pro 등 3개 상용 LLM의 평가결과를 직종별로 병렬 수집한 뒤 양상블 평균으로 통합하였다.

사회적 압력은 BigKinds 중앙일간지 뉴스자료를 기반으로, 직종별 기사수와 논조를 정량화하여 구성하였다. 기사량은 특정 직종의 AI 대체 이슈가 공론장에서 얼마나 자주 호출되는지를 나타내는 의제 강도의 근사치로, 논조는 해당 이슈가 혁신/효율의 프레임으로 수용되는지 혹은 위협/불안·갈등의 프레임으로 확산되는지에 대한 규범적·정서적 방향성의 근사치로 해석하였다. 본 연구는 사회적 압력을 단일 변수로 단순화하지 않고, 기사량과 부정 논조(부정%)를 함께 고려하여 관심(노출)과 갈등(부정성)이 동시에 커지는 경우 정책결정자에게 가해지는 압력(규제 요구)이 강화된다는 점을 반영하였다.

더하여 <표 13>의 RB:RS:CR 권장비율은 연구모형에서 제시한 세 가지 매핑 원칙—(1) 안전·책임 위험 → RB, (2) 운영 중 위험 관찰 → RS, (3) 사회적 수용성 위험 → CR—에 따라 산정된 값이다. 동일 위험유형에 속하는 직종은 동일 비율을 부여하고(A유형 25:35:40, B유형 45:40:15, C유형 15:55:30, D유형 30:50:20), 이는 위험유형이

규제수단 조합을 결정하는 1차 단위라는 본 연구의 설계 원칙을 반영한다. 직종 간 미세한 위험도 차이를 권장비율에 반영하지 않은 것은 본 연구가 직종별 맞춤형 비율 산정이 아니라 유형 단위의 일관된 규제 포트폴리오 설계 원칙을 제안하는 데 목적이 있기 때문이다.

〈표 13〉 평가결과 종합 및 위험도 유형화 - 규제 권장 비율 연계

직종	표준화(평균)	확률화(평균)	AI 대체율(평균(%))	기술위험도(LLM합의)	기사수	긍정/부정/중립(%)	사회적 압력	위험도 유형	RB:RS:CR (권장비율)
법무사	8	6.7	78.5	높음	27	14.8 / 22.2 / 63.0	중간	D	30:50:20
관세사	8	7	75.7	높음	27	44.4 / 14.8 / 40.7	낮음	D	30:50:20
세무사	7	7	69.9	높음	126	35.7 / 11.1 / 53.2	중간	D	30:50:20
공인 회계사	7	6.7	68	높음	47	31.9 / 17.0 / 51.1	중간	D	30:50:20
감정 평가사	6.7	6.7	65.5	높음	42	26.2 / 7.1 / 66.7	낮음	D	30:50:20
의료 기사	6.7	7	68.1	높음	평가 제외				
약사	7	7.3	71	높음	179	35.2 / 11.2 / 53.6	높음	B	45:40:15
공인 노무사	7	6.7	66.9	높음	67	35.8 / 23.9 / 40.3	높음	B	45:40:15
변호사	5.3	6	56.6	중간	2,531	23.2 / 16.0 / 60.8	높음	A	25:35:40
의사	5.3	6.3	53.6	중간	5,840	29.1 /	높음	A	25:35:40

						12.5 / 58.4			
간호사	5	5.3	46.6	낮음	469	34.8 / 10.7 / 54.6	높음	A	25:35:40
변리사	6	6.3	58.8	중간	100	46.0 / 10.0 / 44.0	낮음	C	15:55:30
치과 의사	4.7	5.3	46.8	낮음	69	42.0 / 10.1 / 47.8	낮음	C	15:55:30
한의사	4.3	4.7	42.3	낮음	71	26.8 / 16.9 / 56.3	중간	C	15:55:30
기존 연구와의 대응성		장주희 외(2020)		의사에 대하여 기술대체 가능성을 낮다는 연구결과와 본 연구결과가 일치함					
		송성수(2022)		회계사의 대체확률 = 0.94, 치과외과의 대체확률 = 0.0044로 본 연구 결과와 일치					
		김건우(2018), Frey & Osborne(2017; 2013 워킹페이퍼) 재인용		관세사의 대체확률 = 0.985, 회계사의 대체확률 = 0.957, 세무사의 대체확률 = 0.957, 전문의사의 대체확률 = 0.004, 약사 및 한약사 = 0.012로 대체로 일치하나 약사 및 한약사의 경우 미일치					

구체적으로 <표 13>의 사회적 압력(높음/중간/낮음)은 직종 집합 내에서 기사량의 상대적 규모와 부정 논조의 상대적 강도를 함께 반영해 분류되며, 그 결과 사회적 압력이 높은 직종에는 의사(5,840건), 변호사(2,531건), 간호사(469건) 등 고노출 직종과, 약사(179건), 공인노무사(67건)처럼 노출 규모는 상대적으로 작더라도 부정 논조가 비교적 큰 직종이 포함된다. 반대로 관세사(27건), 법무사(27건) 등 기사량이 낮은 직종은 대체로 사회적 압력이 낮거나 중간 수준으로 분류하였다. 다만 의료기사의 경우 본 단계에서 기사량이 매우 적어 사회적 압력 산정이 불가능하므로, <표 13>에서는 사회적 압력을 NA로 표기하고 위험도 평가에서는 제외 처리하였다.

이상의 두 요인을 결합하여 본 연구는 직종별 위험도를 2×2 구조로 도출하였다. 즉, ① 기술적 대체가능성은 양상블 평균 AI 대체율을 기준으로 구간화하였고(예: 평균 60%

이상을 ‘높음’으로 정의), ② 사회적 압력은 기사량 및 논조 정보를 반영해 ‘높음/중간/낮음’으로 구분하였다. 그 결과 사회적 압력이 높으면서 기술적 대체가능성이 낮은 직종의 유형은 A{고갈등 영역: 사회적 압력 고(高)×기술 대체 저(低)}로, 사회적 압력이 높으면서 기술적 대체가능성이 높은 직종의 유형은 B{고위험 혁신 영역: 사회적 압력 고(高)×기술 대체 고(高)}로 분류된다. 반대로 사회적 압력이 낮거나 중간 수준이면서 기술적 대체가능성이 낮은 직종의 유형은 C{현상 유지 영역: 사회적 압력 저(低)×기술 대체 저(低)}로, 사회적 압력이 낮거나 중간 수준이면서 기술적 대체가능성이 높은 직종의 유형은 D{조용한 대체 영역: 사회적 압력 저(低)×기술 대체 고(高)}로 분류된다.

상기한 <표 13>에서 의사·변호사·간호사는 대체율이 상대적으로 중간 이하(각 53.6%, 56.6%, 46.6%)인 반면 기사량이 매우 높아 사회적 압력이 강하므로 A 유형으로 분류되며, 약사·공인노무사는 대체율이 높고(각 71.0%, 66.9%) 사회적 압력도 높아 B 유형에 위치한다. 변리사·치과의사·한의사는 사회적 압력이 낮거나 중간 수준이면서 대체율이 낮아 C 유형으로 분류된다. 법무사·관세사·세무사·공인회계사·감정평가사 등은 대체율이 높으나 기사량이 낮거나 중간 수준인 경우가 많아 D 유형으로 분류된다(의료기사는 대체율이 높으나 사회 압력 데이터가 없어 제외처리). 해당 결과는 다음 <표 14>로 정리한다.

<표 14> 위험도 유형 및 포함 전문직종

위험도(유형)	정의(사회적 압력 × 기술 대체가능성)	포함 직종
A 고갈등	사회적 압력 높음× 기술 대체가능성 낮음/중간	의사, 변호사, 간호사
B 고위험 혁신	사회적 압력 높음× 기술 대체가능성 높음	약사, 공인노무사
C 현상 유지	사회적 압력 낮음/중간× 기술 대체가능성 낮음	변리사, 치과의사, 한의사
D 조용한 대체	사회적 압력 낮음/중간× 기술 대체가능성 높음	법무사, 관세사, 세무사, 공인회계사, 감정평가사

〈표 13〉의 기존 연구와의 대응성 항목과 〈표 14〉의 유형은 본 연구의 LLM 기반 대체율 판정이 선행 전문가 평가와 어떻게 부합하는지를 보여준다. 첫째, 의사의 경우 장주희 외(2020)가 의사 5명에 대한 직무조사, 7명에 대한 FGI, 델파이 조사 및 전문가협의회를 통해 AI는 의사 업무의 보조적 역할에 머물고 있으며 완전 대체 수준에는 이르지 못한다고 평가하였는데, 본 연구의 LLM 양상불 결과 역시 의사를 대체율 중간(53.6%)·A 유형(보조적 공존 영역)으로 판정하여 두 평가가 방향성에서 일치한다. 둘째, 송성수(2022)는 회계사의 대체확률을 0.94로, 치과의사의 대체확률을 0.0044로 보고하였는데, 본 연구도 공인회계사를 D유형(고대체, 68.0%), 치과의사를 C유형(저대체, 46.8%)으로 분류해 동일한 방향을 보인다.

셋째, 김건우(2018)가 Frey & Osborne(2017; 2013 워킹페이퍼) 방법론으로 산출한 한국 423개 직업의 대체확률은 관세사(0.985)·회계사(0.957)·세무사(0.957)를 고대체로, 전문의사(0.004)를 저대체로 평가하여 본 연구의 D유형(관세사 75.7%·세무사 69.9%·공인회계사 68.0%) 및 A유형(의사 53.6%) 분류와 방향이 부합한다. 다만 약사·한약사의 경우 김건우(2018)는 0.012의 저대체로 평가한 반면, 본 연구는 약사를 B유형(고대체, 71.0%)으로 판정해 방향이 어긋난다. 이는 김건우(2018)가 재인용한 Frey & Osborne(2017; 2013 워킹페이퍼) 분석이 2013년 시점의 미국 노동시장 자료에 기반하여 약사 업무 전반을 대면 상담·환자 안내 중심으로 평가한 반면, 본 연구는 2026년 시점에서 처방 검토·복약지도 자동화·스마트 약국 도입 등 최근의 기술 환경을 반영한 결과로 해석된다.

종합하면, 비교 가능한 5개 직종(의사·회계사·치과의사·관세사·세무사) 가운데 5개 모두에서 선행 전문가 평가와 방향이 일치하며, 약사를 포함한 6개 중 5개(83.3%)에서 수렴이 확인된다. 이러한 외부 수렴은 Szymanski et al.(2025)이 경고한 전문지식 과업에서의 LLM-as-Judge의 체계적 편향이, 적어도 본 연구에서 독립 인간 전문가 평가와 대조 가능한 직종에 한해서는 관찰되지 않음을 시사한다.

상기한 위험도 유형에 따라 본 연구에서 RB(Rule-Based Regulation)는 시장 진입 이전 단계에서 요구되는 사전적 규제요건(등록, 문서화, 성능검증 계획, 책임주체, 로그·추적성 등)을 의미하며, RS(Reactive Supervision)는 운영 과정에서 민원·사고·성능저하를 모니터링하고 시정명령·제재·중단 등을 통해 교정하는 사후 집행을 의미한다.

CR(Co-Regulation)은 정부·전문직 단체·기업·이용자 대표가 규칙·표준·인증·교육·자율점검 체계를 공동 설계·운영하는 공동규제를 의미한다. 세 수단은 상호배타적 대안이 아니라, 위험도에 따라 비중이 달라지는 ‘규제 포트폴리오’로 기능한다.

즉, 사회적 압력이 높아 갈등 및 신뢰 위험이 큰 경우에는(특히 A 유형) 규제기관 단독의 기술요건 부과만으로는 정책적 정당성과 수용성을 확보하기 어렵기 때문에 CR 비중을 확대하여 이해관계자 협의·정보공개·피해구제 경로를 제도화하고, 동시에 RS를 통해 민원·분쟁·사고에 대한 대응력을 강화한다. 반대로 기술적 대체가능성과 확산 속도가 높아 안전·책임 이슈가 정책 리스크로 증폭될 가능성이 큰 경우(특히 B 유형)에는 사전 요건(RB)을 강화하여 진입 단계에서 최소 안전기준과 책임구조를 확립하되, 실제 운영 데이터에서 드러나는 위험을 통제하기 위해 RS(감사·보고·제재)를 병행하는 구조가 적합하다.

이에 따라, RB-RS-CR의 조합이 규제샌드박스 운영규칙(진입: Entry-운영조건: Conditions-측정 및 평가: Measurement-종료 및 전환: Exit)으로 어떻게 구현되는지를 직종별로 구체화하면 다음과 같다. A 유형(고갈등)에서 진입(Entry)은 대체/자동결정이 아니라 보조·지원 중심 기능에 한정하여 사회적 반발을 완화하고, 운영조건(Conditions)에서는 이용자 고지·설명, 인간전문가의 최종책임, 이의제기·분쟁조정 경로를 의무화해야 한다. 측정 및 평가(Measurement)에서는 오류·안전사고뿐 아니라 민원·분쟁·소송·심판 결과, 서비스 신뢰도 등 갈등·수용성 지표를 핵심 KPI로 설정하여 관리하며, 종료 및 전환(Exit)에서는 안전·신뢰 지표가 일정 수준을 충족할 때 단계적 확대를 허용하고, 반대로 사회적 피해·갈등 누적 시 범위를 축소하거나 중단할 수 있다고 판단된다.

B 유형(고위험 혁신)에서 진입은 사전 등록 및 검증계획 제출을 요구하고, 운영조건에서는 로그·추적성, 인간감독, 책임보험, 데이터 거버넌스 등을 강화해야 할 것으로 보인다. 측정 및 평가에서는 성능 저하·편향·사고 발생을 실시간/주기적으로 보고·감사하고, 종료 및 전환에서는 사전에 설정한 임계치 미달 또는 사고 반복 시 즉각 중단·퇴출을 적용해야 할 것으로 보인다. C 유형(현상 유지)은 과잉규제 대신 RS 중심의 저장도 모니터링과 CR 기반 가이드라인으로 충분하며, 기술 성숙 또는 사회적 관심 변화 시 재분류(A 또는 D로 전환)하는 적응적 메커니즘을 두는 형태가 적절하다고 판단된다.

D 유형(조용한 대체)은 사회적 저관심 속에서도 실제 대체가 확산될 수 있다는 점에서

해당 위험성을 고려하여, 신속 진입을 허용하되 최소 사전 요건(RB-lite)과 감사·로그 모니터링 중심의 RS를 강화해야 할 것으로 보인다. 의료기사처럼 사회적 압력 데이터가 미 확보된 직종은 잠정치로 유형을 부여하되, 향후 기사 데이터 보강을 통해 사회적 압력 등급과 위험도를 재산정하여 그에 따라 RB-RS-CR 비중 및 운영규칙을 조정하는 방식으로 연구모형의 학습성을 유지할 필요성이 있다.

V. 결론 및 정책적 함의

본 연구는 연구질문에 따라 다음과 같은 맥락에서 진행되었다. 이는 ① 전문직 업무가 표준화(규칙·절차로의 번역)와 확률화(데이터·학습문제로의 번역) 조건에서 직종별로 어떻게 달라지며, 그 차이가 AI 대체율로 어떻게 나타나는가, ② 기술적 대체율과 사회적 압력(뉴스 기사량·논조)이 어떤 관계를 갖는가, ③ 두 축의 상호작용을 전제로 할 때 직종별로 차별화된 규제샌드박스 설계 원칙은 무엇인가로 정리된다. 결론적으로, 전문직 AI 규제는 기술적으로 무엇이 가능한가로 설계될 수 없고, 사회적으로 무엇이 갈등·신뢰 리스크를 동반하는가를 함께 고려해야 하며, 이 결합이 곧 규제수단의 포트폴리오(RB-RS-CR)와 샌드박스 운영규칙(Entry-Conditions-Measurement-Exit)을 직종별로 고려해야 하는 수단임을 본 연구의 결과로 확인하였다.

결과를 정리하면, 기술적 대체가능성과 관련하여 3개 LLM 평가를 앙상블 평균으로 통합한 결과, 법무사(78.5%), 관세사(75.7%), 약사(71.0%), 세무사(69.9%), 의료기사(68.1%)가 상위권으로 나타났고, 공인회계사(68.0%), 공인노무사(66.9%), 감정평가사(65.5%)가 뒤를 이었다. 중위권은 변리사(58.8%), 변호사(56.6%), 의사(53.6%), 하위권은 치과의사(46.8%), 간호사(46.6%), 한의사(42.3%)로 도출되었다. 이 분포는 AI 성능에 대한 일반적인 상식이라기보다, 해당 직종의 업무가 서식·절차·규칙으로 재현 가능한지(표준화)와 데이터·피드백 축적을 통해 성능을 측정·개선할 수 있는지(확률화)라는 구조적 조건의 차이가 대체율 격차를 설명한다는 점을 고려하여 기술적 대체 가능성을 분석하였음을 밝힌다.

또한, 같은 맥락에서 표준화·확률화의 직종별 차이는 대체율이 곧바로 공공위험이나

규제강도로 변환되지는 않는다. 예컨대 표준화 평균은 법무사·관세사(각 8.0/10)가 가장 높고, 약사·세무사·공인회계사·공인노무사(각 7.0/10)가 그 다음으로 보고되었으며, 확률화 평균은 약사(7.3/10), 관세사·세무사·의료기사(각 7.0/10)가 상위로 나타났다. 반대로 한의사(표준화 4.3/10, 확률화 4.7/10), 치과의사(표준화 4.7/10, 확률화 5.3/10)는 상대적으로 낮다. 즉, 업무가 규칙·서식화된 영역(표준화)과 데이터·학습문제로 환원되는 영역(확률화)이 겹칠수록 대체율이 높아지는데, 이때 규제 정책은 대체율 자체보다도 ① 어떤 업무 단계에서 ② 어떤 오류가 ③ 어떤 피해를 남는지의 맥락(공익적 피해, 책임소재, 사후구제)까지 결합해야 한다.

다음으로, 사회적 압력과 관련하여 중앙일간지 뉴스 9,595건 분석에서 의사(5,840건)와 변호사(2,531건)가 논의의 대부분을 차지하며 사회적 관심이 압도적으로 집중되는 양상이 확인된다. 감성(긍·부정·중립) 측면에서는 직종별로 대체 공포·갈등 프레임과 혁신·효율 프레임이 교차하는데, 법무사(부정 22.2%)와 노무사(부정 23.9%)는 상대적으로 부정 논조가 높게, 변리사(긍정 46.0%)와 관세사(긍정 44.4%)는 긍정 논조가 비교적 높게 나타난다. 이 결과는 기술적 대체율이 높다는 점이 사회적 논쟁으로 이어질 수 있는 가능성이 높다가 아니라, 생명·권리와 직결된 분야(의료·법률)에서는 기술 수준이 중간이어도 신뢰·책임 이슈로 사회적 압력이 급상승할 수 있음을 보여준다.

세 번째, 규제 설계 원칙의 핵심은 상기한 두 결과를 2×2로 결합한 위험유형화가 규제 정책으로 구체화될 수 있다는 점이다. 본 연구는 ‘사회적 압력 × 기술 대체가능성’의 사분면을 구성하여 A(고갈등: 사회 높음×기술 낮음), B(고위험 혁신: 사회 높음×기술 높음), C(현상 유지: 사회 낮/중×기술 낮음), D(조용한 대체: 사회 낮/중×기술 높음)로 분류하고, 직종을 선별하였다. 구체적으로 A에는 의사·변호사·간호사, B에는 약사·공인노무사, C에는 변리사·치과의사·한의사, D에는 법무사·관세사·세무사·공인회계사·감정평가사(의료기사는 사회압력 데이터 한계로 잠정)가 포함된다. 이 유형 분포는 규제의 초점이 AI의 대체 가능성이 아니라 어디에서 사회적 갈등과 공익피해가 증폭되는가로 이동해야 함을 생각해야하며, 곧바로 규제수단 배합의 논리(포트폴리오)로 연결할 수 있다.

즉, 본 연구에서의 규제 논의는 공익이론, 포획이론과 반응적 규제를 실무적으로 함께 고려해야 한다. 사회적 압력이 큰 영역(A·B)일수록, 규제기관이 기술요건만 부과하면 정당성과 수용성이 취약해질 수 있고, 반대로 직역단체가 소비자 보호를 명분으로 진입장벽

을 강화하는 포획 가능성도 커진다. 따라서 본 연구는 사전 의무(RB: 등록·문서화·검증·설명가능성·인간감독·로그·책임주체), 사후 집행(RS: 사고·민원 보고, 모니터링, 표본감사, 시정·제재·리콜·중단), 공동규제(CR: 정부·지역·기업·이용자 대표의 표준·인증·교육·자율점검 공동 설계)를 위험도에 따라 비중을 달리하는 규제 포트폴리오로 제시하고 이를 유형에 따라 비율로 조정하는 형태로 규제를 구체화 하고자 하였다. 특히, A(고갈 등)에서는 CR 비중을 확대해 이해관계자 협의·정보공개·피해구제 경로를 제도화하고 RS로 갈등·분쟁 대응력을 강화하는 것이 합리적이며, B(고위험 혁신)에서는 RB로 최소 안전기준과 책임구조를 먼저 세운 뒤 RS로 운영 데이터 기반 통제를 병행하는 것이 적합하다고 보인다.

정책 설계로의 구체화는 4단계 샌드박스(Entry-Conditions-Measurement-Exit)로 구성하였다. 진입단계에서는 위험도가 높을수록 사전 심사·등록·검증계획 등 RB를 강화하고, 운영조건에서는 적용 범위(업무 단계·대상·환경)를 제한하며 고지·동의, 전문가(인간) 최종책임, 데이터 거버넌스, 책임보험, 로그·추적성을 조건으로 부과해야 하며, 측정 및 평가단계에서는 오류·안전사고뿐 아니라 A유형에서 특히 중요한 민원·분쟁·소송/심판 결과, 서비스 신뢰도 등 ‘수용성 지표’를 KPI로 관리하고, 종료 및 전환단계에서는 사전 임계치에 따라 제도권 편입(확대)·조건 강화 후 연장·중단/퇴출을 결정함으로써 규제의 학습과 조정을 제도화한다. 이때 D(조용한 대체)는 신속 진입을 허용하되 RB-lite(최소 사전요건)와 RS(감사·로그 모니터링)를 결합해 저관심 속 확산에 대응하고, C(현상 유지)는 과잉규제보다 RS 중심의 저장도 모니터링과 CR 기반 가이드라인으로 충분하되 환경 변화 시 재분류 메커니즘을 두어야 한다.

이상의 결론은 2026년 시행되는 한국 AI 기본법 체계에서 하위 법령(시행령·고시·가이드라인)을 고도화 설계할 때, ‘분야(직종) 특화 위험기준’과 ‘운영데이터 기반 적응규제’를 결합해야 한다는 정책적 함의로 직결된다. 즉, 법률이 제시하는 고영향 AI의 포괄 정의를 그대로 두되, 하위 규범에서 ① 전문직 업무 단계를 기준으로 고영향 판단을 세분화하고, ② 인간 개입 의무화가 필요한 지점을 명확히 지정하며, ③ 검증·로그·책임보험·분쟁조정 등 사후구제 인프라를 샌드박스 조건으로 강제하는 방식이 필요하다. 특히 A유형(의사·변호사·간호사)처럼 사회적 압력이 극단적으로 큰 영역은 ‘대체’가 아니라 ‘보조’ 중심으로 진입을 제한하고, 설명·고지·이의제기·분쟁조정 경로를 조건화해야 정책 정당

성과 현장 수용성을 동시에 확보할 수 있다.

본 연구의 규제는 규제 대상 기술(AI)과 인력전환 정책을 결합할 때 효과가 있을 것으로 보인다. 첫째, 공동규제(CR)를 직역단체+기업의 합의로만 좁히면 포획 위험이 커지므로, 이용자 대표·독립 전문가·공공기관이 참여하는 상설 거버넌스와 이해상충 공개를 제도 요건으로 두어야 한다(특히 A유형). 둘째, AI의 직업 대체와 관련한 측정을 성능지표(정확도)에 머무는 것이 아닌, 민원·분쟁·피해구제·신뢰도 같은 사회적 지표를 동등 KPI로 포함해야 하며, 이는 사회적 압력을 단지 분위기 변수로 두지 않고 규제 운영의 핵심 변수로 전환하는 장치가 될 수 있다. 셋째, D유형에서 대체가 빠르게 확산될 가능성을 고려해, 직무 재설계(사람은 예외·책임·대면 커뮤니케이션으로, AI는 정형 처리로)와 재교육을 샌드박스 참여조건(교육 이수·인증)으로 붙이는 방식도 고려할 만하다.

마지막으로 본 연구에서 제안하는 점은 기술대체 가능성과 사회적 압력의 결합을 규제 설계의 출발점으로 삼아, 직종·업무단계·지역에 따라 조건과 측정을 달리하는 정교한 샌드박스 거버넌스를 구축해야 한다는 점이다. 넷째, 실증 분석에서 직접 도출되지는 않으나, 향후 정책적 확장으로서 지역 단위 차등 적용을 고려할 수 있다. 예컨대 D유형의 정형업무(세무·관세·등기·평가)를 자동화 패스트트랙으로 묶어 신속 실증하되 RB-lite+RS를 유지하고, A유형(의료·법률)은 'AI 사용 고자·전문가 최종판단·이의제기·피해구제' 패키지를 검증하는 이원화가 가능하다. 다만 이러한 '전문행정 자동화 특구'·'지역 직종별 AI 대체 신호등 체계' 등 공간 정책 제안은 본 연구의 분석을 넘어서는 후속 과제를 밝힌다.

상기한 바를 고려함에도 불구하고 본 연구는 다음과 같은 한계점을 가진다. 첫째, LLM을 판정자(LLM-as-a-Judge)로 사용한 설계 자체가 내재적으로 갖는 신뢰도 한계가 존재한다. 논문은 프롬프트·모델·상황에 따라 판정이 흔들릴 수 있음을 전제하고, 평가표 제시·다회 반복 평가·불일치 점검, 그리고 다중 샘플링과 합의 규칙으로 측정오차를 줄이려 하였으나, 타당성에 대한 문제가 있을 수 있다. 다만, 각 직역에 대한 광범위한 조사를 하기 어려운 점, 이전 연구에서 LLM을 판정자로 사용한 점 등에 따라 본 연구는 탐색적인 연구로서의 가치를 가진다고 본다. 또한 기술적 가능성과 현장 대체(채택) 속도 사이의 간극이 클 것으로 예상되며, 본 연구의 대체율은 현 시점 상용 AI로 구현 가능성 등을 포함해 산출되었지만, 실제 대체는 조직의 책임구조, 보험·배상 체계, 고객 신뢰, 규

제 준수비용, 데이터 접근권, 노사관계 같은 비기술 요인에 의해 크게 제약될 수 있다. 다만, 현재 전문직역에 대한 AI 대체 가능성 논의가 지속되고 있는 점에 따라 본 연구의 일부 정책적인 기여가 있을 수 있다고 판단된다.

참고문헌

- Acemoglu, D., & Restrepo, P. (2019). Automation and New Tasks: How Technology Displaces and Reinstates Labor. *Journal of Economic Perspectives*, 33(2), 3-30. <https://doi.org/10.1257/jep.33.2.3>
- Abbott, A. (2014). *The system of professions: An essay on the division of expert labor*. University of Chicago press.
- Autor, D. H. (2015). Why are there still so many jobs? The history and future of workplace automation. *Journal of economic perspectives*, 29(3), 3-30.
- Autor, D. H., Levy, F., & Murnane, R. J. (2003). The Skill Content of Recent Technological Change: An Empirical Exploration. *The Quarterly Journal of Economics*, 118(4), 1279-1333.
- Ayres, I., & Braithwaite, J. (1992). *Responsive regulation: Transcending the deregulation debate*. Oxford University Press, USA.
- Brown, T. B., Mann, B., Ryder, N., et al. (2020). Language Models are Few-Shot Learners. *NeurIPS 2020*.
- Brynjolfsson, E., & Mitchell, T. (2017). What can machine learning do? Workforce implications. *Science*, 358(6370), 1530-1534. <https://doi.org/10.1126/science.aap8062>
- Brynjolfsson, E., Mitchell, T., & Rock, D. (2018). What Can Machines Learn, and What Does It Mean for Occupations and the Economy? *AEA Papers and Proceedings*, 108, 43-47. <https://doi.org/10.1257/pandp.20181019>
- Burrell, J. (2016). How the machine 'thinks': Understanding opacity in machine learning algorithms. *Big data & society*, 3(1), 2053951715622512.
- Cornelli, G., Doerr, S., Gambacorta, L., & Merrouche, O. (2024). Regulatory Sandboxes and Fintech Funding: Evidence from the UK. *Review of Finance*, 28(1), 203-233.
- Eloundou, T., Manning, S., Mishkin, P., & Rock, D. (2024). GPTs are GPTs: Labor

- market impact potential of LLMs. *Science*, 384(6702), 1306-1308.
- Entman, R. M. (1993). Framing: Toward Clarification of a Fractured Paradigm. *Journal of Communication*, 43(4), 51-58. <https://doi.org/10.1111/j.1460-2466.1993.tb01304.x>
- Evetts, J. (2003). The sociological analysis of professionalism: Occupational change in the modern world. *International sociology*, 18(2), 395-415.
- FCA, Regulatory Sandbox (<https://www.fca.org.uk/firms/innovation/regulatory-sandbox>), updated 9 May 2024
- Frey, C. B., & Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254-280. <https://doi.org/10.1016/j.techfore.2016.08.019>
- Gilardi, F., Alizadeh, M., & Kubli, M. (2023). ChatGPT outperforms crowd workers for text-annotation tasks. *Proceedings of the National Academy of Sciences*, 120(30), e2305016120.
- Gu, J., Jiang, X., Shi, Z., Tan, H., Zhai, X., Xu, C., ... & Guo, J. (2024). A survey on llm-as-a-judge. *The Innovation*.
- Koo, T. K., & Li, M. Y. (2016). A guideline of selecting and reporting ICC for reliability research. *J. Chiropractic Medicine*, 15(2), 155-163.
- Kroll, J. A. (2015). *Accountable algorithms* (Doctoral dissertation, Princeton University).
- Leshob, A., Bédard, M., & Mili, H. (2020). Robotic Process Automation and Business Rules: A Perfect Match. *ICETE 2020*, 119-126. <https://doi.org/10.5220/0009886701190126>
- MAS, *FinTech Regulatory Sandbox Guidelines* (November 2016, updated January 2022)
- Posner, R. A. (1974). *Theories of economic regulation* (No. w0041). National Bureau of Economic Research.

- Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act), OJ L 2024/1689
- Stigler, George J. (1971). The theory of economic regulation. *Bell Journal of Economics and Management Science*. 2(1), 3-21
- Szymanski, A., Ziems, N., Eicher-Miller, H. A., Li, T. J. J., Jiang, M., & Metoyer, R. A. (2025, March). Limitations of the llm-as-a-judge approach for evaluating llm outputs in expert knowledge tasks. In *Proceedings of the 30th International Conference on Intelligent User Interfaces* (pp. 952-966).
- Wewerka, J., & Reichert, M. (2020). Robotic Process Automation -- A Systematic Literature Review and Assessment Framework. *arXiv:2012.1195*
- Withey, M., Daft, R. L., & Cooper, W. H. (1983). Measures of Perrow's work unit technology: An empirical assessment and a new scale. *Academy of management journal*, 26(1), 45-63.
- Ytreberg, E. (2002). Erving Goffman as a theorist of the mass media. *Critical Studies in Media Communication*, 19(4), 481-497.
- 권용수. (2022). 규제샌드박스 제도의 의의와 쟁점. *법제*, 696, 45-72.
- 권충훈. (2019). 최근 정권별 '자사고' 관련 언론사 기사의 변화 추이 탐색-빅카인즈 시스템 분석을 통한. *인문사회* 21, 10(6), 1757-1771.
- 김건우. (2018). 인공지능에 의한 일자리 위험 진단. *LG 경제연구원*.
- 김경동, 이시영, & 고길곤. (2020). 텍스트 마이닝을 활용한 고용정책 프레임 연구: 언론기사 및 국회회의록의 비교· 분석을 중심으로. *한국사회와 행정연구*, 30(4), 135-163.
- 김기봉. (2019). 업무 자동화를 위한 RPA 융합 기술 고찰. *융합정보논문지*, 9(7), 8-13.
- 도홍일, 박준오, & 이주락. (2024). 빅카인즈 시스템을 활용한 정보경찰 관련 국내 언론보도 변화 분석. *한국경찰연구*, 23(3), 55-80.
- 박종준. (2020). 한국형 규제 샌드박스 법체계에 관한 소고. *법조*, 69(3), 193-232.
- 박주현, 박현지, 김영범. (2024). 빅카인즈를 활용한 5·18 관련 국내 기사 분석 연구. *정보관*

- 리학회지, 41(1), 107-132.
- 송성수. (2022). 인공지능은 인간의 일자리를 얼마나 대체할 것인가: 인공지능 시대의 기술과 노동에 관한 시론. *코기토*, 9(6), 7-37.
- 오준협, & 고보선. (2022). 빅카인즈 분석을 통한 발달장애인법 시행 후 보도 경향 분석. *한국장애인복지학*, 56(56), 59-90.
- 이민창, & 최성락. (2013). 한국의 규제연구 동향 분석 (1990-2012). *한국사회와 행정연구*, 24(2), 339-366.
- 이용희. (2020). 빅데이터를 활용한 국가정책에 따른 '인공지능' 관련 언론기사 변화과정 분석: 기간별 변화를 중심으로. *한국 IT 정책경영학회 논문지*, 12(6), 2119-2125.
- 장주희, 방혜진, 이윤진, 이진솔, 한상근, & 이승희. (2020). 인공지능 시대의 전문직 직업연구. *한국직업능력개발원*
- 최믿음, & 조은수. (2024). AI 관련 뉴스기사의 보도 프레임 분석. *문화기술의 융합*, 10(6), 339-346.
- 최성희. (2025). 전수교육관의 미디어 동향 분석: 빅카인즈를 중심으로. *한국콘텐츠학회논문지*, 25(5), 651-661.
- 최해옥, & 이광호 (2021). 프로젝트형 규제샌드박스의 이해당사자 간 인식차이 분석을 통한 정책수요 파악에 관한 연구. *규제연구*, 30(1), 55 - 78.

부록

〈각주 5에 대한 부록〉 전문직 별 세부업무 평가항목

변호사	변리사	회계사	세무사	법무사
【법률 문서 검토 및 작성】	【선행기술 조사】	【회계 감사】	【세무신고 대리】	【등기·등록 신청】
계약서 검토 및 리스크 분석	검색전략 수립(키워드·IPC)	감사계획·위험평가 문서화	신고자료(증빙) 수집·점검	신청서 작성
판례 검색 및 법리 적용	특허·논문 DB 검색/수집	감사절차 수행·테스트 기록	신고서 작성(부가/소득/법인)	첨부서류 준비
법령 적합성 검토	요약 및 대비표(차별점) 작성	감사조서·증적 정리	전자신고 제출·오류 보완	등기소 제출 및 접수
최종 의견서 작성	선행기술 조사보고서 작성	감사보고서 초안·Q&A 정리	정정신고·사후 문의 대응	진행 상황 확인
【판례 검색 및 분석】	【명세서 작성】	【재무제표 작성 및 검토】	【세액 계산】	【서류 작성】
판례 데이터베이스 검색	발명 구성·도면 설명 정리	조정분개 검토·정리	과세표준 산정·조정	정관 작성
판례 요지 추출	청구항 초안 작성	재무제표 작성	공제·감면 적용 검토	계약서 작성
판례 경향 분석	명세서 본문(실시예 등) 작성	주식 작성·정합성 점검	세액 산출근거표 작성	위임장·동의서 작성
사안 적용 가능성 판단	보정/보완 문구 작성	공시서류 검토·제출 준비	납부세액 검증·확정	공증서류 준비
【법률 해석 및 자문】	【특허성 판단】	【세무 상담】	【세무조사 대응】	【법률 문서 검토】
법령 조문 해석	신규성 판단·대비표 작성	세무 쟁점 검토	조사 대응자료 준비	등기부등본 분석
사례 분석 및 의견 제시	진보성 판단·근거 정리	세무조정·신고서 검토	소명서·확인서 작성	계약서 법적 검토
클라이언트 상담	기재요건·명확성 검토	세무리스크 진단	질의응답 정리·기록	하자 및 리스크 파악
전략적 법률 조언	거절이유 대응논리 작성	유권해석/질의서 작성	불복절차 서면 초안	개선 의견 제시
【소송 전략 수립】	【특허 전략 수립】	【경영 자문】	【세무 상담】	【법률 상담】
사실관계 정리 및 쟁점 도출	출원 포트폴리오·일정	원가·수익성 분석	세무기초(장부/증빙) 상담	등기 절차 안내

	수립			
증거 분석 및 평가	우선권·분할·PCT 전략	예산·사업계획 수립 지원	상속·증여 상담	권리관계 설명
소송 시나리오 설계	특허맵·회피설계(FTO) 분석	재무모델·가치평가(DCF)	거래유형 세무처리 안내	분쟁 예방 조언
승소 가능성 평가	라이선스·기술이전 자료 작성	KPI·경영보고서 작성	사전답변/질의서 작성	사후 관리 안내
【법정 변론 및 협상】	【심판·소송 대리】	【내부통제 평가】	【절세 전략 수립】	【소송 서류 작성】
변론 준비서면 작성	불복심판 청구서·이유서 작성	프로세스 문서화(Flow/RCM)	법인전환·급여/배당 설계	소장·답변서 작성
법정 구두변론	무효/정정심판 서면 작성	통제설계 평가	비용처리·감가상각 최적화	증거자료 정리
상대방과의 협상	증거자료 정리·제출	운영테스트·증적 수집	가업승계·증여 전략 검토	준비서면 작성
즉석 법률 대응	구술심리/준비서면 준비	개선권고·평가보고서 작성	연간 절세 실행계획 수립	소송 진행 보조
공인노무사	관세사	감정평가사	의사	치과의사
【4대 보험 업무】	【품목분류(HS코드)】	【시세 조사】	【영상 판독】	【구강 영상 판독】
취득·상실·변경 신고(EDI)	제품 사양 분석·기준 검토	비교사례 수집·선별	영상 소견 요약·우선순위 분류	파노라마/CT 소견 요약
보수충액 신고/정산	HS 후보코드 검색·비교	공시/실거래 DB 확인	이상소견 탐지·근거 정리	치주·치근단 병변 평가
보험료 산정·납부 관리	분류 근거(조항/사례) 정리	가격수준·추세 분석	판독리포트 초안 작성	소견 표준입력·리포트
산재·고용보험 신청서 작성	사전심사/의견서 작성	시세조사 표·요약 작성	임상정보 통합 해석	치료계획·영상 근거 매칭
【급여 계산】	【관세 계산】	【데이터 분석】	【검사 결과 해석】	【충치·질환 진단】
근태데이터 정리·매핑	과세가격 산정(운임/보험)	가격요인 변수화·정리	이상치 탐지·추세 분석	구강검진 기록·정리
연장·야간·휴일수당 계산	관세·부가세 세액 계산	기초모형 적용·보정	추가검사/재검 판단 근거	우식·치주 진단·분류
연차·퇴직금 산정·검증	관세율·감면·FTA 적용 검토	보정계수·민감도 검토	환자용 결과 설명 요약	통증 원인 감별진단 정리
급여명세서 작성·검토	정산·사후조정 검증 리포트	표·그래프 작성·해석	경과기록·모니터링 정리	소견서·보험코드 입력
【취업규칙 작성】	【수출입 신고】	【현장 조사】	【진단 및 처방】	【치료 계획 수립】
법령 적합성 진단	서류	물적현황 확인·기록	감별진단 목록화	치료옵션

	검토(Invoice/BL 등)			비교(보존/보철/교정)
취업규칙·인사규정 초안	전자신고(EDI) 작성·수정	주변환경·접근성 조사	가이드 기반 처방·오더세트	비용·기간 산정·계획서
의견청취·신고 절차 준비	첨부 제출·통관 진행 관리	규제·권리관계 확인	금기·상호작용·용량 체크	치료 순서·일정 설계
개정 공지·교육자료 작성	검사/보완요구 이슈 대응	현장조사서·증빙 정리	처방전·진단서·기록 작성	동의서·설명자료 준비
【노동 상담】	【FTA 활용 자문】	【평가서 작성】	【환자 상담】	【치료 시술】
계약·임금체불 상담	원산지 기준 판단	평가방법 선정·근거 구성	질환교육·생활습관 상담	보존치료(수복 등) 수행
징계·해고 절차 상담	원산지 서류 준비·검토	요인비교표 작성·보정	치료옵션 설명·동의서 안내	근관치료 수행·기록
산재·휴가·휴업 상담	사후검증 대응자료 작성	평가액 산정·검산	위험·부작용 Q&A 지원	임플란트/외과 처치 수행
자문서·답변서 작성	FTA 매뉴얼·교육자료 작성	평가서 편집·제출	상담기록·추적계획 수립	시술기록·보험청구 입력
【노사 분쟁 조정】	【관세 심판·소송】	【전문가 판단】	【수술·시술】	【환자 상담】
사실관계 정리·쟁점화	처분 분석·쟁점 정리	특수물건/권리제한 판단	수술 전 평가·계획	구강위생·예방 교육
합의/조정안 초안	이의/심사/심판청구 서면	시장·리스크 반영 조정	수술/시술 기록·코딩	치료 후 관리 안내
노동위 구제신청 서면	증빙·감정자료 정리·제출	질의응답·설명자료 작성	수술 중 변수 대응 판단	통증·불만 대응
조정회의 준비·이행관리	의견서·준비서면 작성	이의신청 대응 의견서	수술 후 합병증 모니터링	재내원 계획·기록
한의사	약사	간호사	의료기사	
【맥진·설진】	【처방 검토(DUR)】	【활력징후 측정】	【임상병리검사】	
문진·초진기록 입력	환자정보·처방전 확인	활력징후 측정·EMR 입력	검체 접수·전처리·추적	
설진·맥진 소견 기록	DUR 경고 확인·분류	이상치 감지·보고	분석검사 수행(자동/수동)	
증상 점수화·지표화	의사 문의·조정 기록	추세 모니터링·요약	QC·이상치 재검(델타체크)	
추적관찰 항목 설정	조제 전 최종검증	기기 점검·교정 기록	결과 입력·보고	
【변증】	【약물 상호작용】	【투약 관리】	【방사선 검사】	

	확인]		
변증유형 후보 도출	병용약·OTC·건기식 확인	투약 처방 확인·준비	금기 확인·환자 준비
증상·변증 매칭 검토	상호작용/부작용 위험평가	투약 수행·기록	촬영 프로토콜 설정·촬영
동반질환·금기 고려	대체약·모니터링 제안	반응 관찰·부작용 모니터링	영상 품질 점검·재촬영 판단
변증결론·근거 기록	약력기록 업데이트	투약오류 보고·예방	PACS 업로드·기록
【처방 결정】	【복약지도】	【환자 관찰 및 기록】	【검사 결과 분석】
방제/약재 구성안 작성	복용법·기간·시간 안내	간호사정 기록	이상치 판정·결과 검증
용량·가감·복용법 설계	부작용·주의사항 교육	간호중재 수행·기록	논리검증·해석 지원
주의사항·상호작용 점검	순응도 확인·상담기록	I/O·통증·상처 기록	추가검사/재검 권고 기록
처방전 발행·기록	라벨·안내문 제공	인제서·기록 요약	결과 요약 전달
【침구 시술】	【조제】	【정서적 돌봄】	【검사 장비 관리】
혈위 선택·시술계획	피킹·계수·혼합 조제	심리상태 스크리닝·상담	장비 점검·캘리브레이션
시술기록 표준화 작성	포장·라벨링	보호자 교육·소통	시약/소모품 재고·유효기간
반응 모니터링·경과기록	조제기록·재고 반영	퇴원교육 자료 제공	오류로그·A/S 요청
부작용·응급대응 기록	조제오류 이증검수	민원·갈등 초기대응 기록	정기 성능검사 기록
【환자 상담】	【약료 상담】	【응급 대응】	【환자 안전 관리】
섭생·생활습관 상담	만성질환 약물관리(MTM)	응급 인지·코드 호출	환자 식별·검사확인
치료경과·추적계획	예방접종·OTC 상담	CPR/응급처치 수행	피폭/차폐 등 안전수칙
복약·시술 후 주의 안내	의약품정보 제공·문서화	응급기록·사건보고	감염·검체 안전관리
예약·추적관리	치료성과 모니터링·소통	사후 디브리핑·개선 정리	사고보고·예방교육

A Study on Designing a Regulatory Sandbox for Professional-Occupation Substitutability by Artificial Intelligence (AI/LLMs) and Social-Pressure Dynamics

Han Byeol Yoo, Soung Wan Kim

This study examines how advances in generative artificial intelligence (AI) affect professional occupations (e.g., physicians and lawyers) through a dual lens of technical substitutability and social pressure, and proposes a Korean-style operational model for a regulatory sandbox aligned with these dynamics. Using national daily newspaper coverage from 2022 to 2026 to proxy social pressure, we combine this evidence with task-level substitution assessments generated via large language models (LLMs)—namely GPT, Claude, and Gemini—to assess technical substitutability across 14 professional occupations and estimate risk profiles for occupations with sufficient social-pressure data. To mitigate single-model dependence, we ensemble the assessments of the three models and validate inter-model measurement reliability using the intraclass correlation coefficient (ICC), Krippendorff's α , and Kendall's W. The results indicate that judicial scriveners and customs brokers exhibit high levels of technical substitutability yet receive relatively limited social attention. By contrast, physicians and lawyers fall into a high-conflict domain characterized by moderate technical substitutability but

exceptionally high levels of social pressure. Drawing on capture theory and responsive regulation, we develop a four-stage regulatory sandbox tailored to occupational characteristics. In particular, we propose a differentiated strategy that (i) mandates human-in-the-loop intervention in high-risk domains while (ii) establishing an automation fast track for low-risk administrative tasks. Finally, the study derives policy implications for designing subordinate regulations under the Korean AI Basic Act, which came into effect in 2026.

Key Words: Artificial intelligence, Professional Occupations, Standardization, Probabilization, Regulatory Sandbox, Social Pressure, Media Framing