

법무법인 종사자의 생성형 인공지능 수용 과정에 관한 연구

권 선 영†

한남대학교 문헌정보학과

본 연구는 법무법인 종사자들의 생성형 인공지능(Generative AI) 수용 과정을 탐색하고, 법률 업무 맥락에서 나타나는 인식과 경험의 특성을 분석하는 데 목적이 있다. 이에 본 연구는 기술수용모형(TAM)을 분석틀로 활용하여 법무법인 종사자들의 지각된 유용성, 사용 용이성, 신뢰 인식, 위험 인식 및 조직 맥락을 질적으로 분석하였다. 연구는 국내 법무법인 종사자 7명을 대상으로 반구조화 인터뷰를 실시하여 자료를 수집하였다. 분석 결과, 참여자들은 생성형 AI를 문서 정리, 초안 작성 및 아이디어 도출을 지원하는 보조적 도구로 인식하고 있었다. 반면 허위 정보 생성, 책임 귀속의 불명확성, 기밀 정보 유출 가능성 등에 대한 우려 역시 강하게 나타났다. 또한 생성형 AI 수용은 개인의 기술 친숙성만으로 결정되기보다 조직의 책임 구조와 내부 규범, 검증 요구와 밀접하게 연결되어 있었다. 본 연구는 법률서비스 분야에서 생성형 AI 수용이 조건부·제한적 형태로 나타나고 있음을 보여주며, 향후 법무법인 차원의 가이드라인과 조직적 대응 체계 마련의 필요성을 시사한다.

주요어: 생성형 인공지능, 법무법인, 기술수용모형, 법률서비스, 질적 연구

† 단독저자 : 권선영, 한남대학교 문헌정보학과 교수, 대전광역시 대덕구 한남로 70, E-mail : sykw@hnu.kr
■ 초초투고일 : 2026년 5월 18일 ■ 심사마감일 : 2026년 6월 22일 ■ 게재확정일 : 2026년 7월 15일

1. 서론

최근 몇 년 사이, 법무법인 내부의 업무 환경에서도 눈에 띄는 변화가 감지되고 있다. 판례 검색, 자료 정리, 문서 초안 작성과 같이 반복적이고 시간 소요가 큰 업무 과정에서 생성형 인공지능(Generative AI)이 하나의 보조 도구로 거론되기 시작한 것이다. 과거에는 법률 리서치와 문서 작성이 전적으로 개인의 경험과 숙련도에 의존해 이루어졌다면, 이제는 인공지능이 제시하는 초안이나 요약물을 참고하는 방식이 점차 일상적인 선택지로 등장하고 있다. 이러한 변화는 단순히 새로운 도구의 도입이라기보다, 법률업무 수행 방식 전반에 대한 인식의 전환을 요구하는 흐름으로 볼 수 있다.

생성형 인공지능은 대규모 텍스트 데이터를 기반으로 자연어 결과를 생성·요약·분석할 수 있다는 점에서, 문서 중심 업무 비중이 높은 법률 실무와 결합 가능성이 큰 기술로 논의된다. 실제로 업계 보고서 및 실무 논의에서는 생성형 AI가 법률조사, 문서 요약, 계약·의견서 초안 작성 및 검토 등에서 활용될 수 있으며, 업무 생산성과 처리 속도를 개선할 잠재력을 갖는 것으로 제시된다(Thomson Reuters, 2025; Kemp, 2025). 또한 법률 실무·제도 맥락을 다루는 연구들은 생성형 AI가 법률 텍스트를 생성하거나 요약하는 과정에서 논점을 구조화하고 문서 작성 과정을 보조하는 방식으로 법조 업무의 일부 절차를 변화시킬 수 있음을 논의한다(Regalia, 2024; Contini, 2024).

그러나 이러한 기대와 동시에, 법률실무 현장에서는 생성형 AI 활용에 대한 우려와 회의적 시각이 지속적으로 제기되고 있다. 무엇보다 생성형 AI의 특성상 오류(환각)나 부정확한 결과가 발생할 수 있으며, 법률문서·법률조사에서의 오기재

나 허구 인용은 곧바로 전문직 윤리와 책임 문제로 연결된다. 국내에서도 생성형 AI 활용이 법률서비스에 가져올 영향과 함께 정확성·책임·윤리 쟁점이 핵심 과제로 논의되어 왔다(정채연, 2024; 이창규, 2024). 해외에서도 AI가 생성한 ‘존재하지 않는 판례’나 ‘부정확한 인용’을 법원 제출 문서에 포함한 사례가 반복적으로 보고되며, 법원 제재 및 책임 논쟁이 부각되고 있다(Thomson Reuters, 2025; AP News, 2025).

이러한 맥락에서 생성형 AI의 법률업무 활용 여부는 ‘기술을 들여오느냐’의 문제가 아니라, 종사자의 인식과 판단, 그리고 조직의 규범·통제 방식과 결합된 수용 과정으로 이해될 필요가 있다. 즉, 생성형 AI가 실제로 도움이 된다고 인식되는지(유용성), 사용이 어렵거나 번거롭지 않은지(용이성), 결과를 어느 수준까지 신뢰할 수 있는지(신뢰), 그리고 오류·책임·기밀·데이터 보호 위험을 어떻게 평가하는지(위험 인식)가 수용 태도와 의도를 형성하는 핵심 요소로 작동한다. 법조 윤리와 기밀 유지 의무를 전제로 한 가이드라인들도 이러한 점을 전면에 두고, 정확성 검증·비밀유지·전문성(역량) 확보·거버넌스 체계 마련을 반복적으로 강조한다(Law Society of Ireland, 2025; Terzidou, 2025; CCBE, 2025). 다시 말해, 법무법인에서의 생성형 AI 활용은 기술 자체의 기능보다도 이를 둘러싼 인식 구조와 조직적 관리(가이드라인, 책임 배분, 검증 절차)의 문제에 가깝다.

그럼에도 불구하고 기존의 생성형 AI 및 기술 수용 관련 연구들은 주로 일반 기업, 공공기관, 교육 및 의료 분야를 중심으로 축적되어 왔으며, 법무법인 종사자(변호사 및 지원 스태프 포함)를 대상으로 한 연구는 상대적으로 제한적인 실정이다. 업계 보고서들은 법률 분야의 생성형 AI 도입 확산과 활용 영역을 제시하지만, 실제 법률조직 내

부에서 수용이 어떤 논리와 감정, 규범과 책임 구조 속에서 전개되는지에 대한 질적 설명은 상대적으로 약하다(Thomson Reuters, 2025). 최근 법률전문가 대상 조사나 비교 보고서가 등장하고 있으나, 상당수는 도입 수준·사용 실태·태도 분포를 파악하는 데 초점이 맞추어져 있어, 조직 문화·위계·윤리 규범·리스크 관리가 수용 과정에 미치는 영향을 심층적으로 드러내는 데에는 한계가 있다(LexisNexis, 2023). 특히 법무법인은 위계적 구조와 강한 직업윤리, 책임 중심의 업무 특성을 갖는 조직으로, 동일한 기술이라도 수용 과정이 다른 조직과 다른 양상으로 전개될 가능성이 크다(Terzidou, 2025).

이에 본 연구는 법무법인 종사자를 대상으로 생성형 AI 수용 요인을 질적 관점에서 탐색하고자 한다. 기술수용모형(Technology Acceptance Model; TAM)을 이론적 분석틀로 삼고 심층 인터뷰를 통해 생성형 AI에 대한 지각된 유용성, 지각된 사용 용이성, 신뢰 및 위험 인식이 수용 태도와 수용 의도 형성 과정에서 어떠한 방식으로 작동하는지를 맥락적으로 분석한다. 또한 법률 실무에서 핵심으로 제기되는 윤리·기밀·정확성 문제와 조직 차원의 거버넌스 요구를 함께 고려함으로써, 법률업무 맥락에서 생성형 AI 수용 과정의 복합적 특성을 입체적으로 조명하고, 향후 법무법인 차원의 기술 도입 전략 및 내부 가이드라인 논의에 기초 자료를 제공하고자 한다(Kemp, 2025; Law Society of Ireland, 2025).

본 연구의 연구 문제는 다음과 같다.

- 연구문제 1. 법무법인 종사자는 생성형 AI의 법률업무 활용 가능성을 어떻게 인식하고 있는가?
- 연구문제 2. 생성형 AI에 대한 지각된 유용성과 지각된 사용 용이성은 법무법인 종사자

의 수용 태도 형성에 어떠한 방식으로 작용하는가?

- 연구문제 3. 법무법인 종사자는 생성형 AI의 신뢰성과 위험을 어떻게 인식하며, 이러한 인식은 수용 또는 유보 판단에 어떤 영향을 미치는가?
- 연구문제 4. 법무법인이라는 조직적·전문직 맥락은 생성형 AI 수용 과정에 어떠한 영향을 미치는가?

2. 연구의 방법

본 연구는 법무법인 종사자의 생성형 인공지능(Generative AI) 수용 요인을 심층적으로 이해하기 위해 질적 연구 방법을 채택하였다. 생성형 AI의 법률업무 활용은 단순한 기술적 효용의 문제가 아니라, 전문직 윤리, 책임성, 기밀성, 조직 문화 등 복합적인 맥락 속에서 인식되고 판단되는 현상이므로, 설문조사를 통한 계량적 분석만으로는 그 수용 과정을 충분히 설명하기 어렵다. 이에 본 연구는 법무법인 종사자의 실제 경험과 인식을 중심으로 수용 태도 형성 과정을 탐색하고자 반구조화된 심층 인터뷰(semi-structured interviews)를 주요 자료 수집 방법으로 활용하였다.

이론적 분석틀로는 기술수용모형(Technology Acceptance Model, TAM)을 적용하였다. TAM은 새로운 정보기술에 대한 사용자의 수용 과정을 설명하는 대표적인 이론으로, 지각된 유용성(perceived usefulness)과 지각된 사용 용이성(perceived ease of use)이 기술에 대한 태도와 수용 의도에 영향을 미친다고 설명한다(Davis, 1989). 본 연구는 TAM을 변수 간 인과관계를 검증하기 위한 통계적 모형으로 사용하기보다는 생성형 AI 수용과 관련된 핵심 인식 요인을 체계적

으로 해석하기 위한 분석틀(analytic framework)로 활용하였다.

최근 생성형 AI 연구에서는 결과의 신뢰성, 책임성, 윤리, 위험과 같은 요소가 기술 활용에 중요한 영향을 미치는 요인으로 논의되고 있다(Dwivedi et al., 2023). 특히 법률서비스 분야는 전문직 윤리와 책임성, 기밀 유지 의무가 강하게 요구되는 영역이라는 점에서 이러한 요인들이 생성형 AI 수용을 이해하는 데 중요한 의미를 가진다. 이에 본 연구에서는 TAM의 기본 구성요인에 신뢰와 위험 인식을 함께 고려하여 생성형 AI에 대한 지각된 유용성, 지각된 사용 용이성, 신뢰, 위험 인식이 법무법인 종사자의 수용 태도 및 수용 의도 형성 과정에서 어떠한 의미를 갖는지를 질적으로 분석하였다.

1) 자료 수집

본 연구의 대상은 국내 법무법인에 근무하는 변호사 및 법률 사무직 종사자이다. 질적 연구의 특성을 고려하여 통계적 대표성보다는 정보 풍부성(information-rich cases)을 기준으로 표집을 진행하였다. 인터뷰 참여자는 생성형 AI에 대한 인식이나 사용 경험의 유무, 직무 유형(변호사/사무직), 경력 연차 등을 고려하여 목적 표집(purposive sampling) 방식으로 선정하였다(<표 1> 참조).

<표 1> 참여자 인적사항

참여자	직무	성별	경력 연수	법무법인 규모
A	변호사	남	20년 이상	중소형
B	변호사	여	10년 내외	중소형
C	변호사	여	5년 내외	중소형
D	사무직	여	20년 이상	중소형
E	사무직	남	10년 내외	중소형
F	사무직	여	7년 내외	중소형
G	사무직	여	3년 내외	중소형

총 인터뷰 참여자는 7명으로, 변호사 3명과 사무직 종사자 4명으로 구성되었다. 참여자들은 모두 동일한 중소형 법무법인에 재직하고 있었으며, 동일한 조직 환경에서 생성형 AI 활용 경험과 인식을 공유하고 있었다. 인터뷰는 연구 주제와 관련된 주요 인식 범주에서 응답 내용이 반복적으로 나타나고, 후반부 인터뷰에서 새로운 핵심 주제나 의미 있는 범주가 추가적으로 도출되지 않는 상태에 도달했다고 판단되는 시점에서 종료하였다. 특히 생성형 AI에 대한 보조적 인식, 제한적 신뢰, 책임 및 기밀성 우려, 조직 차원의 가이드라인 필요성과 같은 주제가 참여자 유형과 관계없이 반복적으로 나타났다.

자료 수집은 2026년 1월 13일 진행되었으며, 참여자 1인당 약 30분 내외의 심층 인터뷰를 실시하였다. 인터뷰는 대면 또는 비대면 방식으로 진행되었고, 모든 인터뷰는 참여자의 동의를 얻어 녹음 후 전사(transcription)하였다.

인터뷰는 반구조화된 질문지를 기반으로 하여, 생성형 AI에 대한 전반적 인식, 법률업무 활용 경험, 기대와 우려, 수용 또는 유보 판단의 이유 등을 중심으로 진행되었다. 질문은 기술수용모형의 주요 개념을 반영하되, 참여자의 응답에 따라 추가 질문을 통해 경험과 인식을 보다 구체적으로 탐색하였다.

수집된 인터뷰 자료는 주제분석(thematic analysis) 방법을 통해 분석하였다. 분석 과정은 다음과 같은 단계로 이루어졌다.

첫째, 전사된 인터뷰 자료를 반복적으로 읽으며 전체적인 맥락과 주요 의미를 파악하였다.

둘째, 의미 단위에 따라 개방 코딩(open coding)을 실시하여 생성형 AI 수용과 관련된 핵심 발화를 식별하였다.

셋째, 개방코딩을 통해 도출된 코드를 유사한 의미와 맥락에 따라 통합하여 하위 범주를 구성하였

으며, 이를 기술수용모형(TAM)의 주요 구성요소와 비교·해석하여 상위 범주로 체계화하였다(<표 2> 참조).

넷째, 범주 간의 관계를 비교·해석하여 법무법인 종사자의 생성형 AI 수용 과정에서 나타나는 주요 주제와 맥락적 특성을 도출하고, 이를 기술수용모형의 분석틀과 비교·해석하였다.

이 과정에서 단일 발화의 의미를 단편적으로 해석하기보다는, 발화가 등장한 맥락과 참여자의 직무 특성을 함께 고려하여 분석의 신뢰성을 높이고자 하였다.

본 연구는 연구 참여자의 자발적 참여를 전제로 하였으며, 인터뷰에 앞서 연구 목적과 자료 활용 방안을 충분히 설명하였다. 모든 참여자는 익명으로 처리되었고, 인터뷰 자료는 연구 목적 외의 용도로 사용되지 않았다. 또한 연구 결과 제시 과정에서 개인이나 특정 법무법인을 식별할 수 있는 정보는 모두 제거하였다.

질적 연구의 신뢰성과 타당성을 확보하기 위해 다음과 같은 방법을 적용하였다.

첫째, 인터뷰 자료의 반복 검토와 체계적 코딩을 통해 분석의 일관성을 유지하였다.

둘째, 주요 분석 결과는 선행연구 및 기존 논의

와 비교·대조하여 해석의 타당성을 점검하였다. 또한 인터뷰 자료의 반복 검토와 범주 간 비교를 통해 해석의 일관성을 유지하고자 하였다. 분석 과정에서는 참여자의 직무 유형 및 응답 맥락을 함께 고려하여 단일 사례에 과도하게 의존하지 않도록 하였다.

셋째, 필요시 인터뷰 발화를 인용하여 분석 결과가 실제 인터뷰 자료에 근거하고 있음을 명확히 제시하였다.

2) 분석틀

본 연구는 기술수용모형(Technology Acceptance Model, TAM)을 이론적 분석틀로 활용하여 법무법인 종사자의 생성형 AI 수용 과정을 질적으로 해석하였다. TAM은 새로운 정보기술에 대한 사용자의 수용 과정을 설명하는 대표적인 이론으로, 지각된 유용성과 지각된 사용 용이성이 기술에 대한 태도와 수용 의도에 영향을 미친다고 설명한다(Davis, 1989). 본 연구에서는 다양한 기술수용모형 가운데 TAM을 선택한 이유는 생성형 AI 수용 요인 간의 인과관계를 검증하기보다, 인터뷰 자료를 체계적으로 분류하고 해석하기 위

<표 2> 인터뷰 자료의 상위, 하위 범주

상위 범주	하위 범주
생성형 AI에 대한 전반적 인식	보조도구 인식, 업무환경 변화, 인간 중심 판단
활용 경험	판례 탐색, 문서 초안 작성, 자료 요약, 문서 정리
지각된 유용성	업무 효율성, 시간 절약, 반복업무 감소, 논점 정리
지각된 사용 용이성	쉬운 접근성, 직관적 인터페이스, 프롬프트 수정, 검증 부담
신뢰 인식	제한적 신뢰, 결과 검증, 업무별 신뢰 차이
위험 인식	책임 부담, 기밀정보 보호, 윤리 문제, 증거 신뢰성
수용 태도	조건부 수용, 보조도구 활용, 인간 최종 판단
수용 의도	활용 확대, 조건부 사용, 조직 승인 필요
조직 맥락	조직 분위기, 상급자 영향, 가이드라인 부재
제도·가이드라인	활용 기준, 교육, 보안체계, 검증 절차

한 분석틀로 활용하는 데 적합하다고 판단하였기 때문이다. TAM은 핵심 인식 요인을 중심으로 참여자의 경험과 인식을 구조적으로 해석할 수 있다는 점에서 질적 연구의 분석틀로도 활용되고 있다.

기존 TAM이 지각된 유용성(perceived usefulness)과 지각된 사용 용이성(perceived ease of use)을 중심으로 기술수용을 설명하였다면, 본 연구는 법률업무의 특수성을 반영하여 신뢰(trust)와 위험인식(perceived risk)을 핵심 해석 범주로 추가하였다. 최근 생성형 AI 연구에서는 결과의 신뢰성, 책임성, 윤리 및 위험 요인이 기술 활용에 중요한 영향을 미치는 요소로 논의되고 있으며(Dwivedi et al., 2023), 특히 법률서비스 분야는 전문직 윤리와 책임성, 기밀 유지 의무가 강하게 요구되는 영역이라는 점에서 이러한 요인들이 생성형 AI 수용을 이해하는 데 중요한 의미를 가진다.

법무법인은 전문직 윤리와 책임성이 강하게 작

동하는 조직으로, 생성형 AI 수용 여부는 기술적 효용뿐 아니라 정확성, 책임 소재, 기밀성, 윤리적 적합성에 대한 인식과 밀접하게 연결되어 있다. 이에 본 연구는 TAM의 핵심 구성요소와 법률업무 맥락에서 중요하게 제기되는 인식 요인을 결합한 분석틀을 설정하였다.

본 연구에서는 기술수용모형(TAM)을 변수 간 인과관계를 검증하기 위한 연구모형이 아니라, 인터뷰 질문 구성과 자료 해석을 위한 이론적 준거(theoretical framework)로 활용하였다. 인터뷰 질문은 TAM의 핵심 구성요소인 지각된 유용성, 지각된 사용 용이성을 중심으로 구성하였다. 생성형 AI의 특성을 고려하여 신뢰와 위험 인식에 관한 질문을 추가하였으며, 법률서비스 분야의 전문직 특성과 책임성을 반영하기 위해 조직 맥락과 제도적 요구에 관한 질문도 함께 포함하였다(<표 3> 참조).

<표 3> 인터뷰 분석 범주 및 주요 질문

분석 범주	주요 질문 내용
생성형 AI에 대한 전반적 인식	생성형 AI에 대해 처음 인지하게 된 계기와 전반적인 인식은 어떠한가?
활용 경험	법률업무와 관련하여 생성형 AI를 사용해 본 경험이 있는가? 있다면 어떤 업무에서 활용하였는가?
지각된 유용성	생성형 AI가 법률업무에 실질적으로 도움이 된다고 느낀 경험이 있는가? 그렇다면 어떤 측면에서 그러한가?
지각된 사용 용이성	생성형 AI의 사용 과정은 전반적으로 쉽다고 느끼는가, 혹은 부담스럽다고 느끼는가? 그 이유는 무엇인가?
신뢰 인식	생성형 AI가 제공하는 정보나 결과를 어느 수준까지 신뢰할 수 있다고 생각하는가?
위험 인식	생성형 AI 활용과 관련하여 가장 우려되는 점은 무엇인가? (예: 오류, 책임, 기밀성, 윤리 등)
수용 태도	생성형 AI 활용에 대해 전반적으로 긍정적·중립적·부정적 태도 중 어느 쪽에 가까운가?
수용 의도	향후 법무법인 차원에서 생성형 AI가 도입된다면 개인적으로 활용할 의향이 있는가?
조직 맥락	법무법인의 조직 문화나 내부 규범이 생성형 AI 활용에 어떤 영향을 미친다고 보는가?
제도·가이드라인	생성형 AI 활용을 위해 조직 차원에서 가장 필요하다고 생각되는 조건은 무엇인가?

인터뷰는 반구조화(semi-structured) 방식으로 진행하였으며, 참여자가 실제 경험과 인식을 자유롭게 진술할 수 있도록 설계하였다. 수집된 자료는 개방코딩을 통해 의미 단위를 도출한 후 유사한 의미를 중심으로 하위 범주를 구성하였으며, 이를 기술수용모형(TAM)의 주요 구성요소와 비교·해석하여 생성형 AI 수용 과정을 이해하고자 하였다(<표 2> 참조).

3. 연구 결과

본 연구는 인터뷰 내용을 바탕으로 생성형 AI에 대한 전반적 인식, 활용 경험, 지각된 유용성 및 사용 용이성, 신뢰와 위험 인식, 수용 태도와 수용 의도, 조직 맥락 및 제도·가이드라인 요구 등을 중심으로 분석을 수행하였다.

1) 생성형 AI에 대한 전반적 인식

인터뷰 결과, 법무법인 종사자들은 생성형 AI를 법률업무 전반을 대체하는 기술이라기보다는, 반복적이고 시간 소요가 큰 업무를 보조하는 실무 도구로 인식하고 있었다. 참여자들은 생성형 AI를 통해 자료 검색, 문서 요약, 초안 작성 및 아이디어 정리 등이 가능하다고 평가하였으며, 전반적으로 “업무를 보다 편하게 만들어 주는 기술”이라는 인식을 공유하고 있었다.

그러나 생성형 AI에 대한 인식은 기술적 이해 수준에 따라 차이를 보였다. 일부 참여자는 생성형 AI의 작동 원리나 구조를 명확하게 이해하지 못한 상태에서도 업무 과정에서 부분적으로 활용하고 있었으며, 생성형 AI를 “검색 엔진의 연장선” 또는 “참고용 도구” 정도로 인식하고 있었다. 반면 일부 참여자는 생성형 AI가 향후 법률서비

스 환경 자체를 변화시킬 가능성이 있다고 평가하기도 하였다.

특히 참여자들은 생성형 AI를 독립적 판단 주체로 인식하기보다는, 인간 전문가의 판단을 보조하는 수단으로 이해하는 경향을 보였다.

“처음에는 단순히 유행처럼 생각했는데, 실제로 써보니 업무 속도를 줄여주는 부분은 분명 있더군요.” (참여자 A)

“주변에서 많이 사용하길래 자연스럽게 접하게 됐어요. 검색보다 편하다는 느낌이 들었습니다.” (참여자 B)

“처음 사용할 때도 크게 어렵지는 않았어요. 질문만 잘 입력하면 원하는 형태의 답을 비교적 쉽게 얻을 수 있었습니다.” (참여자 F)

일부 참여자는 생성형 AI를 “내비게이션 같은 존재”라고 표현하며, 방향 제시나 아이디어 제공 측면에서는 유용하지만, 실제 판단과 책임은 결국 인간이 수행해야 한다고 설명하였다. 이는 법무법인 종사자들의 생성형 AI 인식이 기술 낙관론이나 기술 거부 의 이분법적 구조가 아니라, 실무적 효용과 제한 가능성을 동시에 고려하는 형태로 나타나고 있음을 보여준다.

2) 생성형 AI 활용 경험

참여자들의 생성형 AI 활용 경험은 직무 유형과 업무 성격에 따라 다소 차이를 보였다. 변호사 참여자들은 주로 판례 탐색, 법률 논점 정리, 문서 초안 작성 및 아이디어 도출 과정에서 생성형 AI를 활용하고 있었으며, 일부는 법률 특화 AI 서비스를 사용해 본 경험이 있다고 응답하였다. 특히

생성형 AI는 사건의 초기 검토 단계에서 다양한 논리 구조나 쟁점을 빠르게 정리하는 용도로 활용되고 있었으며, 장문의 자료를 요약하거나 문서 초안을 구성하는 과정에서도 일정 부분 사용 경험에 나타났다.

다만 변호사 참여자들의 활용 방식은 대부분 ‘보조적 활용’ 수준에 머무르고 있었다. 참여자들은 생성형 AI를 최종 법률 판단이나 공식 문서 작성 단계에 직접 활용하기보다는, 자료 탐색이나 사고 정리 단계에서 제한적으로 사용하는 경향을 보였다. 특히 생성형 AI가 제시하는 판례나 법률 정보에 대해서는 반복적인 검토가 필요하다고 인식하고 있었으며, 결과를 그대로 활용하기보다는 참고 자료 수준에서 활용하는 경우가 많았다.

반면 법률 사무직 종사자들은 문서 정리, 요약, 자료 분류 및 텍스트 정돈과 같은 행정·지원 업무 중심으로 생성형 AI를 활용하고 있었다. 특히 반복적인 자료 정리나 문서 검토 과정에서 생성형 AI가 일정 수준 업무 부담을 줄여준다고 평가하는 경우가 나타났다.

“사건 내용을 처음 정리할 때 큰 흐름을 잡는 용도로 사용해본 적이 있습니다.” (참여자 C)

“초안은 참고할 수 있지만 그대로 제출하기에는 아직 불안합니다.” (참여자 G)

또한 생성형 AI 활용은 대부분 개인적·비공식적 차원에서 이루어지고 있었다. 참여자들은 조직 차원에서 생성형 AI 활용이 공식적으로 장려되거나 체계적으로 도입된 상태는 아니라고 설명하였으며, 실제 활용 여부 역시 개인의 판단과 필요성에 따라 달라지는 경향을 보였다. 일부 참여자는 조직 내부에서 생성형 AI 활용과 관련된 명확한 기준이나 가이드라인이 마련되어 있지 않기 때문

에 적극적으로 활용하기 어렵다고 응답하기도 하였다.

반면 생성형 AI를 실제 업무에 거의 활용하지 않고 있는 경우도 있었다. 그러나 이러한 경우에도 참여자들은 생성형 AI 자체에 대해 전면적으로 부정적인 태도를 보이기보다는, 정확성 문제나 책임 부담 때문에 신중한 태도를 유지하고 있다고 설명하였다. 특히 법률서비스 분야에서는 결과의 오류가 법적 책임 문제와 직접 연결될 수 있다는 점에서, 생성형 AI 활용 자체가 조심스러운 문제로 인식되고 있었다.

종합하면, 참여자들의 생성형 AI 활용 경험은 직무 유형에 따라 일정한 차이를 보이고 있었으나, 전반적으로는 문서 정리, 자료 탐색 및 초안 작성과 같은 보조적·반복적 업무를 중심으로 제한적으로 활용되고 있는 것으로 나타났다. 이는 생성형 AI가 법률서비스 분야에서 아직 독립적 판단 도구라기보다, 업무 효율성을 보완하는 지원 기술의 형태로 활용되고 있음을 보여준다.

3) 지각된 유용성

인터뷰 결과, 대부분의 참여자는 생성형 AI가 법률업무 수행 과정에서 일정 수준의 실질적 유용성을 가진다고 인식하고 있었다. 특히 참여자들은 생성형 AI가 단순히 시간을 절약하는 기능적 도구를 넘어, 업무 과정에서의 심리적 부담을 줄이고 사고의 출발점을 제공해 준다는 점에서 의미가 있다고 평가하였다. 일부 참여자는 생성형 AI를 활용할 경우 막연한 상태에서 업무를 시작하는 부담이 줄어들며, 초기 방향 설정이나 논리 구조 정리에 도움을 받을 수 있다고 설명하였다. 이는 생성형 AI의 유용성이 단순한 작업 속도 향상뿐 아니라, 업무 수행 과정에서의 인지적 부담 감소와도 연결되어 있음을 보여준다.

“예전에는 비슷한 안내문이나 정리 문서를 계속 손으로 수정했는데, 지금은 초안을 먼저 만들어주니까 시간이 훨씬 줄었습니다.” (참여자 E)

또한 참여자들은 생성형 AI가 반복적이고 소모적인 업무를 일정 부분 경감시켜 준다고 인식하고 있었다. 특히 자료 탐색이나 문서 검토 과정에서 핵심 내용을 빠르게 파악할 수 있다는 점은 업무 효율성과 집중도 향상에 긍정적인 영향을 미치는 요소로 평가되었다. 일부 참여자는 생성형 AI 활용 이후 단순 반복 업무에 투입되는 시간이 감소함에 따라, 보다 핵심적인 검토나 판단 업무에 집중할 수 있게 되었다고 설명하였다. 이러한 인식은 생성형 AI가 법률서비스 분야에서 인간 전문가의 역할 자체를 대체하기보다, 전문적 판단 이전 단계의 업무 부담을 완화하는 방향으로 유용성이 형성되고 있음을 시사한다.

한편 참여자들은 생성형 AI의 유용성이 업무 유형에 따라 다르게 나타난다고 인식하고 있었다. 비교적 구조화된 문서 정리, 자료 요약, 표현 정돈 등의 업무에서는 유용성이 높게 평가된 반면, 법률적 해석이나 전략적 판단과 같이 맥락적 사고와 전문적 책임이 요구되는 영역에서는 유용성이 제한적이라는 인식이 나타났다. 특히 실제 사건의 특수성이나 세부적 법률 맥락까지 충분히 반영하지는 못한다고 설명하기도 하였다. 이는 법률서비스 분야에서 생성형 AI의 유용성이 업무 특성과 위험 수준에 따라 선택적으로 인식되고 있음을 보여준다.

또한 참여자들은 생성형 AI 활용 과정에서 결과 자체보다도 ‘생각의 출발점’을 제공해 준다는 점에 의미를 부여하는 경우가 많았다. 일부는 생성형 AI를 통해 새로운 표현 방식이나 논리 전개 방식을 참고할 수 있다고 설명하였으며, 이는 생

성형 AI가 단순 정보 제공 도구를 넘어 사고 보조 도구로 기능할 가능성을 시사한다. 특히 초안 단계에서 다양한 방향성을 제시해 준다는 점은 실무 과정에서 일정 수준 긍정적으로 평가되고 있었다.

그러나 참여자들이 인식하는 유용성은 절대적이거나 독립적인 형태는 아니었다. 다수의 참여자는 생성형 AI가 업무를 완전히 대체할 수 있다고 보지는 않았으며, 최종적인 판단과 검토는 결국 인간 전문가가 수행해야 한다고 인식하고 있었다. 일부 참여자는 생성형 AI에 대해 “없으면 불편할 수는 있지만, 없어도 업무 자체가 불가능한 것은 아니다”라고 설명하기도 하였다. 이는 생성형 AI가 법률서비스 분야에서 필수적 핵심 기술이라기보다, 업무 효율성을 보완하는 선택적 도구로 인식되고 있음을 보여준다.

“자료를 빠르게 요약해주는 부분은 상당히 편리하다고 느꼈습니다.” (참여자 B)

종합하면, 법무법인 종사자들이 인식하는 생성형 AI의 유용성은 단순한 시간 절약이나 업무 자동화 차원을 넘어, 반복 업무 부담 완화, 사고 보조, 초기 방향 설정 지원 등 업무 과정 전반의 효율성을 높여주는 측면에서 형성되고 있었다. 동시에 이러한 유용성은 법률적 판단과 책임이 요구되는 핵심 영역보다는, 보조적·지원적 업무 영역에서 더욱 강하게 인식되는 경향을 보였다.

4) 지각된 사용 용이성

인터뷰 결과, 대부분의 참여자는 생성형 AI 자체의 사용 방법은 비교적 어렵지 않다고 인식하고 있었다. 참여자들은 일반적인 대화형 인터페이스가 직관적이며, 별도의 전문 기술이나 프로그래

밍 지식 없이도 기본적인 활용이 가능하다고 설명하였다. 특히 기존 검색 서비스와 유사한 방식으로 질문을 입력하면 결과를 얻을 수 있다는 점에서, 초기 접근 자체에 대한 심리적 부담은 크지 않은 것으로 나타났다. 일부 참여자는 “생각보다 사용 방법은 단순하다”거나 “일반 검색보다 대화 형에 가까워 오히려 편하다”고 평가하기도 하였다. 이는 생성형 AI가 기술적 진입 장벽 측면에서는 비교적 낮은 수준의 접근성을 가지고 있음을 보여준다.

특히 문서 요약, 자료 정리, 문장 표현 수정 및 초안 생성과 같은 단순 활용 수준에서는 사용 과정 자체가 어렵지 않다고 인식하는 경향이 나타났다. 참여자들은 비교적 짧은 명령어나 질문만으로도 일정 수준의 결과를 얻을 수 있다는 점에서 편리함을 느끼고 있었으며, 일부는 기존 검색 방식보다 결과가 직관적으로 제시된다는 점을 긍정적으로 평가하였다. 또한 생성형 AI가 긴 문장을 자연스럽게 정리하거나 다양한 표현 방식을 제안해 준다는 점에서, 문서 기반 업무가 많은 법률서비스 환경과 일정 부분 잘 맞는다는 인식도 나타났다.

그러나 참여자들이 경험하는 어려움은 단순한 기술 조작의 문제가 아니라, 법률업무 특유의 정확성과 맥락성을 어떻게 반영할 것인가의 문제와 밀접하게 연결되어 있었다. 참여자들은 생성형 AI가 일반적인 수준의 결과는 비교적 쉽게 생성하지만, 실제 법률업무에서 요구되는 세부적 맥락이나 표현의 엄밀성을 충분히 반영하지 못하는 경우가 많다고 인식하고 있었다. 이에 따라 원하는 결과를 얻기 위해서는 질문 내용을 반복적으로 수정하거나, 세부 조건을 구체적으로 입력해야 하는 경우가 빈번하다고 설명하였다.

또한 생성형 AI 활용 과정에서 ‘질문을 어떻게 입력하느냐’ 자체가 중요한 작업이 된다고도 보고

있었는데, 원하는 결과를 얻기 위해 반복적으로 표현을 조정하거나 맥락을 추가해야 한다는 것이다. 이러한 과정은 일반적인 정보 검색과는 다른 형태의 피로감으로 이어지기도 하였다. 특히 법률 업무와 같이 표현의 정확성과 논리적 구조가 중요한 분야에서는 질문 설계 자체가 업무의 일부처럼 인식되는 경향이 나타났다.

“질문을 어떻게 하느냐에 따라 결과 차이가 커서, 익숙해지는 과정은 필요했습니다.” (참여자 B)

또한 참여자들은 생성형 AI 결과를 검토하고 수정하는 과정에서 추가적인 부담을 경험하고 있었다. 생성형 AI가 생성한 문장 자체는 자연스럽더라도, 실제 법률적 의미나 판례 맥락까지 정확하게 반영되어 있는지는 별도의 검토가 필요하다고 인식하고 있었기 때문이다. 일부 참여자는 결과를 반복적으로 확인하고 수정하는 과정에서 오히려 업무 흐름이 길어지는 경우도 있다고 설명하였다. 이는 생성형 AI의 사용 용이성이 단순히 ‘쉽게 사용할 수 있는가’의 문제가 아니라, 결과를 신뢰 가능한 수준으로 다듬기 위해 어느 정도의 추가 작업이 필요한가와도 연결되어 있음을 보여준다.

종합하면, 참여자들은 생성형 AI의 기술적 접근성과 기본 사용 방법 자체는 비교적 용이하다고 평가하고 있었으나, 실제 법률업무에 적용하는 과정에서는 질문 설계, 결과 검토 및 맥락 보완과 같은 추가적 작업 부담을 함께 경험하고 있었다. 이는 법률서비스 분야에서 생성형 AI의 사용 용이성이 단순한 인터페이스 편의성 차원을 넘어, 전문직 업무 특유의 정확성과 책임 요구 속에서 복합적으로 형성되고 있음을 시사한다.

5) 신뢰 인식

참여자들은 생성형 AI가 제공하는 정보와 결과에 대해 전반적으로 제한적 신뢰를 부여하는 경향을 보였다. 대부분의 참여자는 생성형 AI를 전적으로 신뢰할 수 있는 도구라기보다는, 참고와 아이디어 도출 수준에서 활용 가능한 보조 수단으로 인식하고 있었다. 특히 법률서비스 분야는 결과의 정확성과 책임성이 매우 중요하게 작동하는 영역이라는 점에서, 생성형 AI 결과를 그대로 수용하기에는 여전히 상당한 불안 요소가 존재한다고 인식하고 있었다. 이러한 인식은 최근 생성형 AI의 이른바 ‘환각(hallucination)’ 문제와 관련된 국내외 논의와도 밀접하게 연결되어 있었다. 실제로 최근 해외에서는 생성형 AI가 존재하지 않는 판례나 허위 법률 인용을 생성하여 법원 제재로 이어진 사례들이 반복적으로 보고되고 있으며, 법률 실무 현장에서도 생성형 AI 결과의 신뢰성과 검증 책임 문제가 중요한 이슈로 부각되고 있다.

참여자들은 생성형 AI가 제시하는 판례, 법률 해석, 사실관계 요약 등에 오류가 포함될 가능성이 있다고 인식하고 있었으며, 이로 인해 생성 결과를 그대로 업무에 활용하는 것에는 상당한 경계심을 드러냈다. 일부 참여자는 “초안이나 방향 설정에는 도움이 되지만 그대로 사용할 수는 없다”거나, “참고는 가능하지만 책임질 수는 없다”고 설명하였다. 이는 생성형 AI가 제공하는 결과가 문장 형태로는 자연스럽고 설득력 있게 제시되더라도, 실제 법률적 정확성이나 사실 관계까지 보장하는 것은 아니라는 인식과 연결되어 있었다. 최근 연구들 역시 생성형 AI 기반 법률 도구들이 일반 챗봇보다 오류를 줄이기는 하였으나, 여전히 일정 수준 이상의 허위 판례 및 부정확한 정보를 생성할 가능성이 존재한다고 지적하고 있다.

또한 참여자들은 생성형 AI의 신뢰 수준이 업무 영역에 따라 다르게 형성된다고 인식하고 있었다. 아이디어 정리, 문서 구조화, 일반적 정보 요약 등에 대해서는 비교적 긍정적인 평가가 나타났으나, 판례 인용이나 법률적 판단, 계약서 작성과 같이 정확성과 책임성이 핵심적으로 요구되는 영역에서는 신뢰 수준이 현저히 낮아지는 경향을 보였다. 이는 생성형 AI에 대한 신뢰가 전면적·일괄적으로 형성되는 것이 아니라, 업무의 위험 수준과 결과 책임 정도에 따라 선택적으로 구성되고 있음을 보여준다. 일부 참여자는 생성형 AI가 반복적이고 일반적인 정보 처리에는 일정 수준 유용할 수 있으나, 실제 사건의 특수성과 맥락까지 충분히 반영하기에는 한계가 있다고 설명하였다.

“그렇듯하게 틀린 답을 주는 경우가 있어서 그대로 신뢰하기는 어렵습니다.” (참여자 A)

“처음에는 믿고 사용했다가 나중에 다시 확인하는 경우가 많았습니다.” (참여자 D)

특히 참여자들은 생성형 AI의 신뢰 여부가 기술 자체의 성능만으로 결정되는 것이 아니라, 이를 사용하는 사람의 검토 능력과 전문성에 의해 상당 부분 좌우된다고 인식하고 있었다. 즉, 생성형 AI 결과를 그대로 받아들이기보다, 사용자가 법률적 지식과 경험을 바탕으로 결과를 검증하고 수정할 수 있어야 실제 활용이 가능하다는 것이다. 이러한 인식은 최근 법률 AI 연구에서 제기되는 ‘과잉 의존(overreliance)’ 문제와도 연결된다. 일부 연구는 생성형 AI가 자연스럽고 권위적인 형태의 문장을 생성하기 때문에, 사용자가 결과를 충분히 검증하지 않을 경우 오히려 잘못된 판단으로 이어질 가능성이 있다고 지적하고 있다.

또한 일부 국가와 법원은 생성형 AI 활용에 대한 별도의 검증 원칙이나 윤리 기준을 마련하려는 움직임을 보이고 있으며, 이는 법률서비스 분야에서 생성형 AI 활용 시 인간 전문가의 검토 책임이 더욱 중요해지고 있음을 보여준다.

종합하면, 참여자들의 생성형 AI 신뢰 인식은 단순한 기술 신뢰 여부의 문제가 아니라, 법률서비스 분야의 책임 구조와 전문직 윤리, 그리고 결과 검증 가능성과 깊게 연결되어 있었다. 즉, 법무법인 종사자들은 생성형 AI를 일정 수준 활용 가능한 기술로 평가하면서도, 실제 신뢰는 인간 전문가의 검토와 책임 수행을 전제로 한 제한적·조건적 형태로 형성하고 있는 것으로 나타났다.

6) 위험 인식

참여자들은 생성형 AI 활용과 관련하여 다양한 위험 요소를 인식하고 있었으며, 특히 책임성과 정확성 문제를 가장 중요한 위험으로 언급하였다. 참여자들은 생성형 AI가 생성한 결과가 실제 법률업무 과정에서 충분히 검증되지 않은 상태로 활용될 경우, 법적 분쟁이나 업무상 문제로 이어질 수 있다고 우려하고 있었다. 특히 법률서비스 분야는 작은 오류나 표현 차이 역시 중요한 법적 결과를 초래할 수 있는 영역이라는 점에서, 생성형 AI 활용 과정 자체를 신중하게 접근해야 한다는 인식이 강하게 나타났다.

특히 변호사 참여자들은 법률서비스의 특성상 결과에 대한 책임이 개인과 조직에 직접 귀속된다는 점에서 생성형 AI 활용을 더욱 신중하게 바라보고 있었다. 일부 참여자는 “결국 책임은 사람이 져야 한다”거나 “틀렸을 때 AI가 책임져 주는 것은 아니다”라고 설명하며, 책임 귀속의 불명확성을 중요한 위험 요소로 인식하고 있었다. 이는

생성형 AI 활용 과정에서 발생하는 오류가 단순한 기술적 문제에 그치지 않고, 실제 법률적 책임과 전문직 윤리 문제로 연결될 수 있다는 인식과 관련되어 있었다. 최근 국내외 법률업계에서도 생성형 AI 활용 시 결과 검증 책임이 누구에게 귀속되는지에 대한 논의가 활발하게 이루어지고 있으며, 일부 로펌과 기관은 별도의 내부 검토 절차와 AI 활용 기준을 마련하려는 움직임을 보이고 있다.

“법률 분야는 작은 오류도 책임 문제로 이어질 수 있어서 조심할 수밖에 없습니다.” (참여자 A)

“의뢰인 정보 같은 민감한 내용을 입력하는 건 부담스럽게 느껴집니다.” (참여자 D)

또한 기밀 유지와 데이터 보호 문제 역시 주요 위험 요소로 언급되었다. 참여자들은 고객 정보나 사건 관련 자료를 생성형 AI에 입력하는 과정에서 정보 유출 가능성이 존재할 수 있다고 우려하였으며, 이러한 위험 때문에 실제 업무 적용에 제한을 두고 있다고 설명하였다. 특히 법률서비스 분야는 고객 비밀 유지와 민감 정보 보호가 핵심적으로 요구되는 영역이라는 점에서, 외부 AI 서비스 활용 자체에 대한 심리적 부담이 크게 나타났다. 일부 참여자는 “자료를 입력하는 순간 어디까지 저장되는지 알 수 없다”거나 “민감한 사건 내용은 쉽게 넣기 어렵다”고 설명하기도 하였다. 이러한 우려는 최근 국내에서도 생성형 AI 활용 과정에서 개인정보 및 내부 정보 유출 가능성이 중요한 보안 이슈로 논의되고 있는 흐름과 연결된다. 실제 개인정보보호위원회와 정보보호 관련 기관들은 생성형 AI 활용 과정에서 데이터 처리와 내부 거버넌스 구축의 중요성을 지속적으로 강조하고 있다.

한편 일부 참여자는 생성형 AI 기술의 발전이 향후 허위 자료 생성이나 증거 조작 문제로 이어질 가능성에 대해서도 우려를 나타냈다. 특히 음성·영상·텍스트 생성 기술이 빠르게 발전하면서, 향후 법률 실무에서는 디지털 자료의 진위 여부를 검증하는 문제가 더욱 중요해질 수 있다고 인식하고 있었다. 일부 참여자는 “나중에는 진짜와 가짜를 구별하기 더 어려워질 것 같다”거나 “증거 자체를 의심해야 하는 시대가 올 수도 있다”고 설명하기도 하였다. 최근 국내에서도 딥페이크와 생성형 AI 기반 허위 콘텐츠 문제가 사회적 이슈로 확산되면서, 법률·수사·보안 분야에서 AI 생성 자료의 진위 검증 문제에 대한 논의가 증가하고 있다. 과학기술정보통신부 역시 생성형 AI 기반 사이버 위협 증가 가능성을 주요 위협 요소 가운데 하나로 제시하고 있다.

또한 일부 참여자는 생성형 AI 확산이 장기적으로 법률서비스 조직 내부의 업무 구조와 역할 변화로 이어질 가능성에 대해서도 우려를 나타냈다. 특히 반복적 문서 작업이나 자료 정리 업무의 자동화가 확대될 경우, 일부 직무의 역할 축소 또는 업무 재편 가능성이 존재한다고 인식하고 있었다. 이는 생성형 AI 위험 인식이 단순한 기술 오류 수준을 넘어, 향후 조직 운영과 직무 구조 변화에 대한 불안감과도 연결되어 있음을 보여준다.

종합하면, 법무법인 종사자들의 생성형 AI 위험 인식은 단순한 기술 불신 수준에 머무르지 않고, 법률서비스의 책임 구조와 전문직 윤리, 기밀 유지, 데이터 보호 및 조직 변화 가능성 전반과 깊게 연결되어 있었다. 특히 생성형 AI 활용 과정에서 나타나는 위험은 단순한 기능적 문제라기보다, 법률서비스 분야의 핵심 가치인 신뢰성·책임성·윤리성과 직접적으로 연결되는 문제로 인식되고 있는 것으로 나타났다.

7) 수용 태도

인터뷰 결과, 참여자들은 생성형 AI 활용에 대해 전면적인 거부보다는 대체로 조건부 수용 태도를 보이고 있었다. 대부분의 참여자는 생성형 AI가 향후 법률업무 환경에서 일정 수준 이상 활용될 가능성에 대해서는 공감하고 있었으며, 일부는 “이미 흐름이 시작되었다”거나 “앞으로는 사용하지 않을 수 없을 것”이라고 평가하기도 하였다. 특히 반복적이고 시간 소요가 큰 업무의 비중이 높은 법률서비스 환경에서 생성형 AI가 업무 효율성과 생산성을 일정 부분 향상시킬 수 있다는 점에 대해서는 비교적 긍정적인 인식이 공통적으로 나타났다. 일부 참여자는 생성형 AI 활용이 향후 법률서비스 산업 전반의 업무 방식 변화로 이어질 가능성을 언급하기도 하였으며, 젊은 세대를 중심으로 활용 범위가 점차 확대될 것이라는 전망을 제시하기도 하였다.

그러나 이러한 수용 태도는 생성형 AI에 대한 절대적 신뢰나 적극적 의존을 의미하는 것은 아니었다. 참여자들은 생성형 AI가 업무 효율성을 높이는 데 일정 부분 도움이 될 수 있다는 점은 인정하면서도, 최종 판단과 책임은 여전히 인간 전문가가 수행해야 한다고 인식하고 있었다. 특히 법률서비스 분야는 결과의 정확성과 책임성이 매우 강하게 요구되는 영역이라는 점에서, 생성형 AI를 독립적인 판단 주체로 받아들이기에는 여전히 한계가 존재한다고 보았다. 이에 따라 생성형 AI는 법률적 판단을 대체하는 기술이라기보다, 반복적이고 보조적인 업무를 지원하는 수단으로 수용되는 경향을 보였다. 일부 참여자는 생성형 AI를 “참고용 보조 도구” 또는 “초안 작성 수준의 기술”로 인식하고 있었으며, 실제 활용 과정에서도 최종 검토와 수정 과정은 반드시 인간 전문가가 수행해야 한다고 설명하였다.

또한 참여자들의 수용 태도는 업무 영역에 따라 차이를 나타냈다. 자료 탐색, 문서 요약, 초안 작성 및 아이디어 정리와 같은 업무에서는 비교적 긍정적인 태도가 나타났으나, 법률 해석이나 공식 문서 제출, 계약서 작성과 같이 정확성과 책임성이 강하게 요구되는 영역에서는 신중하거나 유보적인 태도가 두드러졌다. 특히 일부 참여자는 생성형 AI가 제시하는 결과의 논리 구조나 표현 방식은 참고할 수 있으나, 판례 인용이나 법률적 판단과 관련된 부분은 반드시 별도의 검증이 필요하다고 인식하고 있었다. 이에 따라 생성형 AI 활용은 단순히 ‘사용 여부’의 문제가 아니라, 어떠한 업무 영역까지 허용 가능한가에 대한 선택적·제한적 판단 과정과 연결되어 있었다.

“완전히 부정적으로 볼 필요는 없다고 생각합니다. 잘 활용하면 도움이 될 수 있습니다.”
(참여자 C)

“실무에서는 이미 어느 정도 필요성이 생기고 있는 것 같습니다.” (참여자 E)

한편 일부 참여자는 생성형 AI 활용 자체보다도, 이를 활용하는 과정에서 발생할 수 있는 책임 문제와 조직 내 시선이 실제 수용 태도에 영향을 미친다고 설명하였다. 즉, 개인적으로는 생성형 AI 활용 가능성을 긍정적으로 평가하더라도, 조직 차원의 기준이나 가이드라인이 충분히 마련되지 않은 상황에서는 적극적으로 활용하기 어렵다는 것이다. 이러한 인식은 생성형 AI 수용이 단순한 개인의 기술 친숙성이나 혁신 성향만으로 결정되는 것이 아니라, 법무법인의 조직 문화와 책임 구조 속에서 형성되는 복합적 과정임을 보여준다.

종합하면, 법무법인 종사자들의 생성형 AI 수용 태도는 단순한 긍정·부정의 이분법적 구조가 아

니라, 업무 특성과 위험 수준, 책임 부담 및 조직 맥락 등을 종합적으로 고려하는 ‘조건부 수용’의 형태로 나타나고 있었다. 이는 법률서비스 분야에서 생성형 AI 수용이 일반적인 정보기술 수용과는 다른 양상으로 전개될 수 있음을 시사한다.

8) 수용 의도

참여자들은 향후 생성형 AI 활용이 법률업무 환경에서 점차 확대될 가능성이 높다고 인식하고 있었으며, 대부분 일정 수준 이상의 활용 의도를 가지고 있는 것으로 나타났다. 특히 반복적이고 행정적인 업무 부담을 줄일 수 있다는 점에서 생성형 AI의 필요성을 긍정적으로 평가하는 경향이 나타났다. 일부 참여자는 생성형 AI가 자료 탐색, 문서 초안 작성, 일정 수준의 정보 정리 업무를 지원하게 될 경우, 법률전문가들이 보다 핵심적인 판단과 전략 수립 업무에 집중할 수 있을 것이라고 전망하였다. 이러한 인식은 최근 국내 법률업계에서도 생성형 AI를 단순한 기술 유행이 아니라, 향후 법률서비스 환경 변화의 주요 요소로 바라보는 흐름과도 연결된다. 실제 국내 로펌과 법률기관들은 생성형 AI 활용 가이드라인과 내부 활용 기준 마련을 논의하고 있으며, AI 활용이 향후 법률서비스 산업 전반에 점진적으로 확산될 가능성이 제기되고 있다.

다만 이러한 수용 의도는 무조건적인 도입 의지라기보다는, 특정 조건이 충족될 경우 활용 가능성이 높아진다는 조건부 형태로 나타났다. 참여자들은 생성형 AI 활용을 위해서는 결과 검증 체계, 책임 기준, 정보 보호 장치 및 조직 차원의 가이드라인이 우선적으로 마련될 필요가 있다고 응답하였다. 일부 참여자는 “기준이 명확해지면 더 적극적으로 사용할 수 있을 것 같다”고 설명하기도 하였다. 특히 참여자들은 생성형 AI 활용 여부 자

체보다도, 문제가 발생했을 때 어떤 기준과 절차에 따라 책임을 판단할 것인지가 더욱 중요하다고 인식하고 있었다. 이는 생성형 AI 수용이 단순히 기술 기능의 우수성만으로 결정되는 것이 아니라, 조직 내부의 통제 체계와 관리 구조에 대한 신뢰와도 밀접하게 연결되어 있음을 보여준다. 최근 정부와 관련 기관 역시 생성형 AI 활용 과정에서 책임성과 투명성, 데이터 관리 체계 마련의 중요성을 강조하고 있으며, 생성형 인공지능 서비스 이용자 보호 가이드라인 등을 통해 조직 차원의 관리 필요성을 제시하고 있다.

“활용 자체는 필요하다고 보지만, 검증 절차는 반드시 같이 가야 한다고 생각합니다.” (참여자 B)

또한 생성형 AI 활용 의도는 조직의 공식 입장이나 내부 분위기와도 밀접하게 연결되어 있었다. 참여자들은 개인적으로는 활용 가능성을 긍정적으로 평가하면서도, 조직 차원의 허용 범위나 책임 구조가 불명확할 경우 실제 활용에는 신중해질 수밖에 없다고 인식하고 있었다. 특히 법무법인과 같은 전문직 조직은 위계 구조와 책임 체계가 강하게 작동하는 환경이라는 점에서, 개인 차원의 기술 활용보다 조직 문화와 내부 규범이 실제 활용 여부에 더 큰 영향을 미친다고 설명하였다. 일부 참여자는 조직 내에서 생성형 AI 활용에 대한 명확한 입장이 정리되지 않은 상황에서는 개인적으로 활용하더라도 이를 공식 업무 과정에 적극 반영하기 어렵다고 응답하기도 하였다. 이는 생성형 AI 수용 의도가 개인의 기술 친숙성만으로 결정되는 것이 아니라, 조직적·제도적 환경과 함께 형성되고 있음을 보여준다.

한편 일부 참여자는 생성형 AI 기술 발전 속도가 매우 빠르기 때문에, 현재는 제한적으로 활용

하고 있더라도 향후에는 보다 적극적인 활용이 불가피할 수 있다고 전망하였다. 특히 젊은 세대를 중심으로 생성형 AI 활용이 자연스럽게 확대될 가능성을 언급하는 경우도 나타났다. 일부는 “지금은 조심스럽지만 결국은 쓰게 될 것 같다”거나 “새로운 세대는 훨씬 자연스럽게 받아들일 것 같다”고 설명하였다. 이러한 인식은 생성형 AI 수용 의도가 현재의 기술 수준에 대한 평가뿐 아니라, 미래 업무 환경 변화에 대한 전망과 기대 속에서도 형성되고 있음을 시사한다. 최근 국내 법조계와 정책 분야에서도 생성형 AI 확산에 대응하기 위한 윤리 기준과 제도적 논의가 활발히 이루어지고 있으며, 이는 생성형 AI 활용이 일시적 현상이 아니라 장기적인 산업 변화 흐름 속에서 이해되고 있음을 보여준다.

9) 조직 맥락

인터뷰 결과, 생성형 AI 수용 과정은 개인의 인식이나 기술 친숙성만으로 설명되기 어려우며, 참여자들이 근무하는 중소형 법무법인의 업무환경과 책임 구조가 생성형 AI 수용에 영향을 미치는 것으로 나타났다. 참여자들은 생성형 AI 활용에 대해 개인적으로는 비교적 긍정적인 태도를 보이더라도, 실제 업무 적용 과정에서는 조직 분위기와 내부 규범을 강하게 의식하고 있었다. 특히 법률서비스 조직은 일반 기업과 달리 전문직 윤리와 책임 체계가 강하게 작동하는 환경이라는 점에서, 생성형 AI 활용 여부 역시 개인의 선택을 넘어 조직 차원의 문제로 인식되고 있었다.

특히 참여자들은 현재 대부분의 법무법인에서 생성형 AI 활용과 관련된 명확한 기준이나 공식 가이드라인이 충분히 마련되어 있지 않다고 인식하고 있었다. 일부 참여자는 “조직 차원에서 명확히 허용하거나 금지한 상태는 아니다”라고 설명

하였으며, 이로 인해 실제 활용은 개인적 판단이나 비공식적 수준에 머무르는 경우가 많다고 응답하였다. 일부 참여자는 내부적으로 생성형 AI 사용 자체를 명시적으로 제한하지는 않지만, 그렇다고 적극적으로 장려하거나 체계적으로 교육하는 분위기 역시 아니라고 설명하였다. 이는 생성형 AI 활용이 공식적 업무 체계 안으로 편입되기 보다는, 개인적 필요와 판단에 따라 제한적으로 이루어지고 있음을 보여준다.

또한 참여자들은 생성형 AI 활용 과정에서 문제가 발생할 경우 책임이 누구에게 귀속되는지가 불명확하다는 점을 중요한 조직적 문제로 인식하고 있었다. 특히 법률서비스는 결과에 대한 책임성과 윤리성이 강하게 요구되는 분야이기 때문에, 생성형 AI 활용 여부 역시 단순한 효율성의 문제가 아니라 조직 차원의 위험 관리 문제와 연결되어 있다고 설명하였다. 일부 참여자는 생성형 AI를 통해 생성된 자료나 문서에서 오류가 발생할 경우, 실질적인 책임은 결국 해당 업무를 담당한 변호사나 조직이 부담하게 될 가능성이 높다고 인식하고 있었다. 이러한 인식은 생성형 AI 활용이 기술적 편의성 차원을 넘어, 조직 내부의 책임 구조 및 검토 체계와 긴밀하게 연결되어 있음을 보여준다.

특히 참여자들이 근무하는 법무법인의 위계 구조 역시 생성형 AI 활용 태도에 일정한 영향을 미치는 것으로 나타났다. 일부 참여자는 조직 내부에서 상급자의 인식이나 조직 분위기가 실제 활용 여부에 상당한 영향을 준다고 설명하였다. 즉, 개인적으로는 생성형 AI 활용에 긍정적이더라도 조직 차원의 분위기나 암묵적 규범 때문에 적극적으로 활용하기 어려운 상황이 존재한다는 것이다. 일부는 “괜히 문제 될 수 있는 일을 먼저 시도하기 어렵다”거나 “상급자가 조심스러운 입장이면 실제 사용도 제한될 수밖에 없다”고 설명하기도

하였다. 이는 생성형 AI 활용이 단순한 기술 선택이 아니라, 조직 내 권한 구조와 문화적 분위기 속에서 조정되는 과정임을 시사한다.

또한 일부 참여자는 법무법인 조직이 전통적으로 보수적이고 안정성을 중시하는 성향을 가지고 있으므로 새로운 기술 도입 자체에 신중한 분위기가 존재할 수 있다고 설명하였다. 특히 법률서비스 분야는 정확성과 신뢰성이 핵심 가치로 작동하는 영역이라는 점에서, 충분한 검증 없이 새로운 기술을 업무에 적용하는 것에 대한 심리적 부담이 상대적으로 크게 나타났다. 최근 국내 법조계에서도 생성형 AI 활용과 관련한 내부 윤리 기준 및 보안 가이드라인 논의가 증가하고 있으나, 실제 조직 차원의 제도화 수준은 아직 초기 단계라는 평가가 제기되고 있다.

“조직에서 명확하게 방향을 정해준 게 없다
보니 직원들도 조심스럽게 사용하는 분위기입니다.” (참여자 E)

한편 일부 참여자는 조직 차원의 명확한 기준과 교육 체계가 마련될 경우 생성형 AI 활용 역시 보다 안정적으로 확대될 수 있을 것이라고 전망하였다. 특히 생성형 AI 활용 범위, 검증 절차, 정보 입력 기준 및 책임 범위 등이 제도적으로 정리된다면 조직 내부에서도 활용에 대한 심리적 부담이 감소할 수 있다고 설명하였다. 이는 생성형 AI 수용이 단순히 개인의 기술 친숙성이나 혁신 성향에 의해 결정되는 것이 아니라, 조직 차원의 관리 체계와 제도적 신뢰 형성과 밀접하게 연결되어 있음을 보여준다.

종합하면, 법무법인에서의 생성형 AI 수용은 개인 단위의 기술 선택이라기보다, 조직 차원의 책임 구조, 윤리 규범, 위험 관리 체계 및 위계적 조직 문화 속에서 형성되는 복합적 과정으로 나타났

다. 특히 참여자들은 생성형 AI 활용 여부 자체보다도, 이를 어떠한 기준과 절차 속에서 관리하고 통제할 것인가를 더욱 중요한 문제로 인식하고 있었다.

10) 제도 및 가이드라인 요구

참여자들은 생성형 AI 활용 확대를 위해 가장 우선적으로 필요한 요소로 조직 차원의 제도 정비와 가이드라인 마련을 언급하였다. 특히 생성형 AI 활용 범위, 검증 절차, 책임 기준 및 정보 보호 원칙 등에 대한 명확한 기준이 필요하다고 인식하고 있었다. 참여자들은 현재 생성형 AI 활용이 대부분 개인의 판단과 경험에 의존하여 이루어지고 있으며, 이로 인해 조직 내부에서도 활용 수준과 방식에 상당한 차이가 존재한다고 설명하였다. 일부 참여자는 “어디까지 활용 가능한지 명확하지 않다”거나 “사용 자체보다 기준이 없는 상태가 더 불안하다”고 언급하기도 하였다. 이는 생성형 AI 활용 과정에서 기술 자체보다도 조직 차원의 관리 체계 부재가 더 큰 불확실성으로 작용하고 있음을 보여준다.

다수의 참여자는 현재 법무법인 내부에서 생성형 AI 활용 여부가 개인의 판단에 상당 부분 의존하고 있다고 설명하였으며, 이로 인해 활용 수준과 방식이 구성원마다 크게 다를 수 있다고 지적하였다. 특히 일부는 동일한 조직 내부에서도 생성형 AI를 적극적으로 활용하는 구성원이 있는 반면, 전혀 활용하지 않거나 의도적으로 회피하는 경우도 존재한다고 설명하였다. 이는 조직 차원의 공식 기준이나 교육 체계가 부재한 상황에서, 생성형 AI 활용이 개별 구성원의 기술 친숙성이나 위험 인식 수준에 따라 달라지고 있음을 시사한다. 일부 참여자는 “최소한 어디까지는 사용 가능하고, 어떤 경우는 제한되는지에 대한 기준이 필

요하다”고 언급하기도 하였다.

“사용 방법이나 주의사항에 대한 기본 교육이 있으면 훨씬 편할 것 같습니다.” (참여자 G)

또한 참여자들은 생성형 AI 활용과 관련된 교육 및 검증 체계의 필요성도 강조하였다. 참여자들은 생성형 AI를 효과적으로 활용하기 위해서는 단순히 기술 사용 방법을 익히는 수준을 넘어, 생성 결과의 오류 가능성을 검토하고 적절히 수정·보완할 수 있는 역량이 중요하다고 인식하고 있었다. 특히 일부 참여자는 생성형 AI 활용 과정에서 “무엇을 입력하면 안 되는가”, “어떤 결과를 그대로 사용하면 위험한가”와 같은 실무적 기준 교육이 필요하다고 설명하였다. 이는 생성형 AI 활용이 단순한 디지털 기술 습득 문제가 아니라, 전문직 업무 수행 과정에서의 판단 역량과도 연결되어 있음을 보여준다.

특히 법률업무의 특성상 고객 정보 보호와 책임 문제가 핵심적으로 작용하기 때문에, 참여자들은 조직 차원의 보안 기준 및 데이터 관리 원칙이 우선적으로 마련되어야 한다고 강조하였다. 일부 참여자는 사건 자료나 고객 정보를 생성형 AI 서비스에 입력하는 행위 자체에 대해 상당한 부담을 느끼고 있었으며, 실제 업무 적용 과정에서도 정보 입력 범위를 스스로 제한하고 있다고 설명하였다. 이러한 우려는 최근 국내외 법률업계에서 생성형 AI 활용 시 데이터 보안 및 개인정보 보호 문제가 주요 쟁점으로 논의되고 있는 흐름과도 연결된다. 실제 국내 법조계와 정보보호 분야에서는 생성형 AI 활용 시 내부 데이터 관리 원칙, 외부 서비스 사용 제한 및 검증 절차 마련의 필요성이 지속적으로 제기되고 있다.

또한 일부 참여자는 향후 법무법인 차원에서 생성형 AI 활용 정책이나 내부 운영 원칙이 마련될

필요가 있다고 전망하였다. 특히 생성형 AI 활용 여부를 개인의 자율 판단에만 맡기기보다는, 조직 차원의 검토 절차와 승인 체계가 함께 구축되어야 한다고 설명하였다. 일부는 향후 법무법인 내부에서도 생성형 AI 활용 전담 조직이나 검토 역할이 필요해질 수 있다고 언급하기도 하였다. 이는 생성형 AI 활용이 단순한 개인 업무 보조 수준을 넘어, 조직 운영 체계 전반과 연결되는 문제로 인식되고 있음을 보여준다.

종합하면, 참여자들은 생성형 AI 활용 가능성 자체에 대해서는 대체로 긍정적으로 평가하고 있었으나, 실제 조직 차원의 활용 확대를 위해서는 책임 구조와 윤리 기준, 보안 체계 및 검증 절차를 포함한 제도적 기반 마련이 선행되어야 한다고 인식하고 있었다. 특히 법률서비스 분야에서 생성형 AI 활용은 단순한 기술 도입 문제가 아니라, 조직의 업무 체계와 전문직 윤리, 위험 관리 구조를 함께 재정비해야 하는 문제로 이해되고 있는 것으로 나타났다.

4. 결 론

본 연구는 법무법인 종사자를 대상으로 생성형 인공지능(Generative AI)의 수용 과정을 질적으로 탐색하고, 기술수용모형(TAM)을 기반으로 지각된 유용성, 지각된 사용 용이성, 신뢰, 위험 인식 및 조직 맥락이 수용 태도와 수용 의도 형성에 어떠한 영향을 미치는지를 분석하였다. 이를 위해 법무법인 종사자 7명을 대상으로 심층 인터뷰를 실시하고, 생성형 AI에 대한 실제 경험과 인식을 중심으로 분석을 수행하였다.

연구 결과, 법무법인 종사자들은 생성형 AI를 법률업무 전반을 대체하는 기술이라기보다 반복적이고 시간 소요가 큰 업무를 보조하는 실무적

도구로 인식하고 있었다. 특히 자료 탐색, 문서 요약, 초안 작성 및 논점 정리와 같은 영역에서는 생성형 AI의 효율성과 활용 가능성을 긍정적으로 평가하였다. 이는 생성형 AI의 지각된 유용성이 업무 처리 속도와 편의성 향상 측면에서 형성되고 있음을 보여준다. 또한 대부분의 참여자는 생성형 AI의 인터페이스 자체는 비교적 직관적이고 접근성이 높다고 평가하였으나, 생성 결과를 검증하고 수정해야 한다는 부담 때문에 실제 사용 과정에서 심리적·업무적 부담을 동시에 경험하고 있는 것으로 나타났다.

반면 생성형 AI에 대한 신뢰 수준은 제한적으로 나타났다. 참여자들은 생성형 AI가 제공하는 결과를 아이디어 도출이나 참고 수준에서는 활용 가능하다고 보았으나, 판례 인용이나 법률적 판단과 같이 정확성과 책임성이 요구되는 영역에서는 신뢰 수준이 현저히 낮아지는 경향을 보였다. 특히 허위 정보 생성, 부정확한 판례 제시 및 책임 귀속의 불명확성은 생성형 AI 활용을 주저하게 만드는 핵심 요인으로 나타났다. 또한 고객 정보 및 사건 자료 입력 과정에서 발생할 수 있는 기밀 유출과 데이터 보호 문제 역시 중요한 위험 요소로 인식되고 있었다.

본 연구는 법무법인 종사자들의 생성형 AI 수용 태도가 단순한 긍정 또는 부정의 이분법적 구조가 아니라, 업무 유형과 위험 수준에 따라 조절되는 ‘조건부 수용(conditional acceptance)’ 형태로 나타난다는 점을 확인하였다. 참여자들은 생성형 AI 활용 가능성 자체에 대해서는 대체로 긍정적 태도를 보였으나, 실제 활용 확대를 위해서는 결과 검증 체계, 책임 기준, 보안 장치 및 조직 차원의 가이드라인이 선행되어야 한다고 인식하고 있었다. 이는 생성형 AI 수용이 개인의 기술 친숙성이나 태도만으로 결정되는 것이 아니라, 법무법인의 조직 문화와 책임 구조, 전문직 윤리 및

내부 규범과 긴밀하게 연결되어 있음을 의미한다.

특히 본 연구는 본 연구 대상인 중소형 법무법인의 사례를 통해 생성형 AI 수용 과정이 일반 조직과 다른 전문직 조직의 특성을 보일 가능성을 확인하였다. 법률서비스는 높은 수준의 정확성과 책임성이 요구되는 분야이며, 결과에 대한 책임이 개인과 조직에 직접 귀속되는 구조를 갖는다. 이에 따라 생성형 AI 활용 여부 역시 단순한 업무 효율성의 문제가 아니라, 전문직 윤리와 위험 관리 체계 속에서 판단되는 경향이 강하게 나타났다. 이는 기존 기술수용 연구들이 충분히 설명하지 못했던 법률전문직 조직 특유의 맥락적 특성을 보여준다는 점에서 의의를 가진다.

이러한 결과는 생성형 AI가 문서 작성과 자료 요약 등에서 업무 효율성을 향상시키는 보조도구로 활용될 수 있다는 기존 연구(Thomson Reuters, 2025; Kemp, 2025)와 대체로 일치한다. 반면 본 연구에서는 법률전문직의 특성상 생성형 AI의 유용성뿐 아니라 신뢰, 책임성 및 위험 인식이 실제 수용 여부를 설명하는 중요한 요인으로 확인되었다. 아울러 이러한 결과는 기존 기술수용 연구를 법률서비스라는 전문직 맥락으로 확장하여 해석하였다는 점에서 의미가 있다.

본 연구의 학문적 의의는 다음과 같다. 첫째, 기존 생성형 AI 수용 연구가 주로 계량적 접근에 집중되어 있던 한계를 보완하여, 질적 연구를 통해 법무법인 종사자의 실제 경험과 인식을 심층적으로 분석하였다는 점이다. 둘째, 기술수용모형(TAM)을 법률업무 맥락에 적용하면서 신뢰와 위험 인식을 핵심 해석 요소로 포함함으로써, 생성형 AI 수용 과정을 보다 입체적으로 설명하고자 하였다. 셋째, 생성형 AI 수용이 개인 수준의 기술 선택을 넘어 조직의 책임 구조와 윤리 규범 속에서 형성되는 과정을 질적으로 제시하였다는 점에서 의의를 가진다.

실무적 측면에서는 법무법인의 생성형 AI 활용 확대를 위해 명확한 내부 가이드라인과 거버넌스 체계 구축이 필요함을 시사한다. 특히 생성형 AI 활용 범위, 검증 절차, 책임 기준, 데이터 보호 원칙 등에 대한 제도적 기준 마련이 중요하며, 단순 기술 교육을 넘어 윤리 및 검증 역량 강화 교육이 함께 이루어질 필요가 있다. 또한 법률전문직의 특성을 고려한 보안 중심 생성형 AI 환경 구축과 조직 차원의 위험 관리 체계 정비 역시 중요한 과제로 제기된다.

한편 본 연구는 일부 한계를 가진다. 우선 본 연구는 동일한 중소형 법무법인에 근무하는 소수의 참여자를 대상으로 수행된 질적 사례연구라는 점에서 연구 결과를 모든 법무법인 종사자 또는 다양한 유형의 법률조직으로 일반화하는 데에는 제한이 있다. 또한 참여자의 직무, 조직 규모 및 생성형 AI 활용 경험 수준에 따라 인식 차이가 존재할 가능성이 있으나, 이를 충분히 비교·분석하지 못하였다. 향후 연구에서는 보다 다양한 규모와 유형의 법률조직을 대상으로 한 비교연구와 함께, 생성형 AI 활용 경험 수준에 따른 수용 차이를 계량적으로 검증할 필요가 있다. 더 나아가 생성형 AI 활용이 실제 법률업무 수행 방식과 전문직 역할 변화에 어떠한 영향을 미치는지에 대한 장기적 연구도 요구된다.

종합하면, 생성형 AI는 법률업무의 효율성을 향상시킬 수 있는 잠재력을 가진 기술로 인식되고 있으나, 본 연구 대상인 중소형 법무법인 종사자들의 생성형 AI 수용은 신뢰·위험·책임 및 조직 규범에 대한 인식과 밀접하게 결합된 조건부 형태로 나타나고 있었다. 따라서 향후 법률 분야에서 생성형 AI의 안정적 활용과 확산을 위해서는 기술 발전 자체뿐 아니라, 전문직 윤리와 책임 구조를 반영한 조직적·제도적 기반 마련이 함께 이루어져야 할 것이다.

참 고 문 헌

- 이창규 (2024). 변호사 업무에서 인공지능 활용과 법률윤리. <법학논총>, 44(2), 77-109.
- 정채연 (2024). 생성형 인공지능을 활용한 법률서비스의 쟁점과 과제. <법학연구>, 35(3), 401-443.
- Associated Press (2025, June 7). UK judge warns of risk to justice after lawyers cited fake AI-generated cases in court. AP News. Available: <https://apnews.com/article/uk-courts-fake-ai-cases-46013a78d78dc869bdfd6b42579411cb>
- Contini, F. (2024). Unboxing generative AI for the legal professions. *International Journal for Court Administration*, 15(1), 1-15. <https://doi.org/10.36745/ijca.604>
- Council of Bars and Law Societies of Europe (CCBE) (2025, October 2). CCBE guide on the use of generative AI by lawyers. Council of Bars and Law Societies of Europe (CCBE). Available: https://www.ccbе.eu/fileadmin/speciality_distribution/public/documents/IT_LAW/ITL_Guides_recommendations/EN_ITL_20251002_CCBE-guide-on-the-use-of-the-use-of-generative-AI-for-lawyers.pdf
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319-340. <https://doi.org/10.2307/249008>
- Dwivedi, Y. K. et al. (2023). So what if ChatGPT wrote it? Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Kemp, R. (2025, March 4). The Rise of Generative AI in Legal Practice. Chambers and Partners. Available: <https://chambers.com/articles/the-rise-of-generative-ai-in-legal-practice>
- Law Society of Ireland (2025, December). Guidelines for the Use of Generative Artificial Intelligence by the Legal Profession in Ireland. Dublin: Law Society of Ireland. <https://www.lawsociety.ie/artificial-intelligence-ai/generative-ai-guidance>
- LexisNexis (2023, April). Generative AI adoption survey among legal professionals. LexisNexis Legal Research Report. https://www.lexisnexis.com/pdf/ln_generative_ai_report.pdf
- Regalia, J. (2024). From Briefs to Bytes: How Generative AI is Transforming Legal Writing and Practice. UNLV William S. Boyd School of Law, Faculty Scholarship. <https://scholars.law.unlv.edu/cgi/viewcontent.cgi?article=2482&context=facpub>
- Terzidou, K. (2025). Generative AI systems in legal practice: Offering quality legal services while upholding legal ethics. *International Journal of Law in Context*, 21(1), 1-18. <https://doi.org/10.1017/S1744552324000065>
- Thomson Reuters (2025). 2025 Generative AI in Professional Services Report. Toronto: Thomson Reuters. Available: <https://www.thomsonreuters.com/en/reports/2025-generative-ai-in-professional-services-report.html>

Exploring Factors Influencing the Acceptance of Generative AI among Law Firm Professionals: A Qualitative Approach Based on the Technology Acceptance Model

Sunyoung Kwon

Hannam University

This study explores the acceptance of generative artificial intelligence (AI) among law firm professionals within the context of legal services. As generative AI technologies such as ChatGPT become increasingly discussed in legal practice, they are being considered for tasks including document drafting, legal research, case summarization, and information organization. However, on the other hand, concerns regarding accuracy, responsibility, confidentiality, and ethical issues remain significant barriers to adoption. This study applies the Technology Acceptance Model (TAM) as an analytical framework and conducts qualitative semi-structured interviews with seven law firm professionals, including attorneys and legal office staff. These findings indicate that the participants generally perceived generative AI as a supportive tool for repetitive and administrative tasks rather than as a substitute for professional legal judgment. The participants acknowledged the usefulness of generative AI in improving efficiency and reducing repetitive workloads, while simultaneously expressing concerns regarding accountability, information reliability, and data protection. The study also found that AI acceptance is strongly influenced by organizational culture, professional ethics, and internal responsibility structures within law firms. These results suggest that generative AI acceptance in legal services is conditional and closely connected to organizational and ethical contexts.

Keywords: Generative AI, Technology Acceptance Model (TAM), Law Firms, Legal Professionals, AI Acceptance, Legal Services, Qualitative Research