

FDA의 새로운 기준과 마주하게 된 AI 신약개발



인공지능(AI)은 신약개발 전 과정에 걸쳐 데이터 생성부터 분석, 규제당국에 제출해야 하는 각종 자료 준비까지 다양하게 활용되고 있다. 제약산업 전반에서 AI 활용은 앞으로 더 늘어날 것으로 전망된다.

이러한 세계적 흐름을 반영하듯, 마틴 마커리 미국 식품의약국(FDA) 국장은 지난 6월 10일 <JAMA>(의학·보건 분야 최상위권 학술지)에 게재된 자신의 기고문을 통해 안전성과 정확성을 유지하면서 혁신을 촉진하기 위해 AI 접근법을 고려하고 AI 기술을 적극 활용하는 것을 FDA의 다섯 가지 핵심 우선순위 중 하나로 제시한 바 있다.

°글 함병균·한지엘 외국변호사(미국) °일러스트 게티이미지뱅크

미국 규제 체계 전반에는 신약개발 분야에서의 인공지능(AI) 활용을 포괄적으로 규율하는 법률 등 AI에 대한 일관적인 연방 차원의 법령이 존재하지 않는다. 이에 따라 규제는 여전히 각 규제 기관에서 발행하는 분야 별 지침을 중심으로 이뤄지고 있으며 의약품 분야의 경우, 미국 식품의약국(FDA)이 일련의 지침을 통해 핵심 규제 방향을 제시한다.



FDA의 지침 초안 발표

FDA는 지난 1월 7일 바이든 행정부 말기이자 트럼프 행정부 출범 전 '의약품 및 생물학적 제품 관련 규제 의사 결정을 지원하기 위한 AI 사용에 관한 고려사항'이라는 제목의 지침 초안(Draft Guidance on Considerations for the Use of Artificial Intelligence to Support Regulatory Decision-Making for Drug and Biological Products, 이하 본 지침)을 발표했다. 이는 FDA가 의약품 및 생물학적 제품 개발에 사용되는 AI에 대한 첫 지침이다. 본 지침은 AI가 의약품 관련 규제 체계에서 어떻게 평가되고 통합되어야 하는지에 대한 FDA의 최신 규제 방향을 반영하는 중대한 이정표다. 본 지침은 특정 사안에 대한 FDA의 현 입장을 기술한 문서이고 법적으로 강제성이 있는 규제 요건이 아닌 권고사항 또는 참고사항으로 간주되어야 한다는 한계가 있다. 하지만 AI를 활용한 의약품 개발에 대해 이를 직접 규율하는 연방 차원의 법령이 부재한 상황에서 여전히 본 지침은 앞으로의 규제 방향성에 대해 파악할 수 있는 중요한 근거로 작용한다.



본 지침이 적용되는 AI 사용 유형

본 지침은 의약품의 안전성, 유효성 또는 품질과 관련된 규제 결정에 활용되는 데이터 또는 정보를 생성하는 다양한 AI 사용 유형들을 포괄적으로 다룬다. 예컨대, 다음과 같은 AI 사용 유형은 본 지침의 적용 대상에 해당할 수 있으나, 이에 국한되지 않는다.

- 임상시험계획(Investigational New Drug) 승인 신청과 관련해 신약개발 과정에서 AI를 사용하는 경우

- 새로운 형태의 임상시험 설계에 AI를 사용하는 경우
- 신약개발 프로그램 내에서 AI 기반 디지털 헬스 기술을 활용하는 경우
- 의약품 제조 공정에 AI를 사용하는 경우
- AI가 활용되는 모델을 기반으로 신약을 개발하는 경우(model-informed drug development)
- 임상 데이터를 분석해 실제 임상 근거를 생성하는 연구에서 AI를 활용하는 경우

다만, AI가 의약품의 후보물질 탐색(drug discovery)이나 제약사의 내부 운영 효율성을 위해 사용되는 경우에는 연구 결과의 신뢰도나 제품의 품질에 영향을 미치지 않는 한 본 지침의 적용 대상에 해당하지 않는다. 예를 들어, 내부조직의 자원 배분, 프로젝트 관리 또는 규제 관련 서류 초안 작성과 같은 업무에 AI가 활용되더라도, 해당 활용이 규제 결정에 사용되는 데이터의 정확성, 타당성 또는 품질에 직접적인 영향을 미치지 않는다면 지침의 적용 대상에 포함되지 않는다.



위험성 기반 신뢰성 평가체계 (Risk-based Credibility Framework)

본 지침은 AI 모델의 출력 결과물이 규제 결정을 돕기 위한 정보 또는 데이터를 생성하는 데 사용되는 경우를 전제로 한다. 이 경우, 모델의 출력 결과에 대한 신뢰성을 입증하기 위해 필요한 정보를 계획, 수집, 정리 및 문서화할 수 있도록 위험성 기반 신뢰성 평가 체계(Risk-based Credibility Framework)를 제시한다. 이 접근 방식은 AI 모델의 특성상 내부 작동 원리를 외부에서 명확히 확인하기 어려운, 이른바 '블랙박스' 구조에서 기인한다. 이로 인해 AI 모델에 대한 신뢰를 확보하기 위해서는 보다 엄격하고 정량적인 증거가 필요하다. FDA는 본 지침에서 신뢰성(credibility)을 "특정 사용 맥락(Context of Use)에서 AI 모델의 성능에 대해 신뢰성을 입증하는 증거를 수집함으로써 형성되는 신뢰(trust)"로 정의하고 있다. 이는 AI 모델에 대한 신뢰는 자동적으로 부여되지 않고 Context of Use에 따라 구체적인 증거로 입증돼야 함을 의미한다. 아울러 본 지침은 단순 기술 성능 외에도 투명성(transparency), 설명 가능성(explainability), 맥락적 적합성(contextual appropriateness) 등 다양한 요소가 AI 모델의 신뢰도



함병균

미국 보건복지부 법무실에서 변호사 경력을 시작해, 이후 저명한 글로벌 로펌들의 제약, 바이오, 헬스케어 자문팀 일원으로 활동했다. 한국에서는 글로벌 분자진단 기업인 시젠에서 사업개발 및 법무 총괄 임원으로 재직하였으며 현재는 덴톤스 리 법률사무소의 생명과학 헬스케어 팀을 이끌며 제약 바이오, 진단, 의료기기, 헬스케어 자문 업무를 수행하고 있다. 제약 바이오 분야 제품 수명주기 전 단계 자문과 기술 라이선싱 및 협업 거래를 중점적으로 다룬다.



한지엘

조지타운 로스쿨을 졸업한 이후 미국 노동부 법무실, 국내외 NGO, 스타트업에서 폭넓은 법률 자문 경험을 쌓았다. 현재는 덴톤스 리 법률사무소에서 생명과학 및 헬스케어 산업을 중심으로 기술 라이선싱, 국제 거래, 규제 대응 등의 분야에 주력하고 있다. 이 외에도 국내외 주요 기업을 대상으로 한 전략적 법률 자문을 제공한다.

~~~~~  
본 지침은  
FDA가 단순히  
산업의 흐름을  
따라가는 것을  
넘어, AI가  
신약개발에  
미치는 영향을  
선제적으로  
규정하고  
이끌고자  
하는 의지를  
보여준다.  
~~~~~

에 영향을 준다는 점을 고려할 것을 권고한다.
본 지침에서 제시하는 위험성 기반 신뢰성 체계는 다음
의 7단계 절차로 구성된다.



7단계 절차

본 지침에서 제시하는 7단계 절차는 Context of Use에
서 사용되는 AI 모델이 신뢰할 수 있는지 평가하고,
FDA가 요구하는 문서화 및 정보 공개에 대비할 수 있도
록 방향성을 제시한다.

1단계

질문(Question of Interest) 정의

가장 먼저, AI 모델이 해소하고자 하는
구체적인 질문 또는 결정 사안을 명확히 정의해야 한
다. 예컨대, 여기에는 “환자의 이상 반응 위험도를 고려
하였을 때 입원 모니터링이 필요한가”와 같은 질문이
해당될 수 있다. AI 모델 사용은 명확한 목적 또는 개발
필요성을 기반으로 해야 하고, 그 목적에 정확하게 부
합해야 한다.

2단계

Context of Use 정의

Context of Use는 AI 모델의 적용 범위,
역할, 해당 모델의 출력 결과가 어떻게 질문에 활용되는
지를 포함해 AI 모델 사용의 구체적 맥락을 설명하는 개
념이다. Context of Use를 정의할 때, 질문에 답하기 위
해서 AI 모델 외에도 다른 종류의 근거 자료가 존재하는
지 여부도 함께 고려되어야 한다. 만약 다른 종류의 근거
자료가 존재하지 않는다면 AI 모델은 해당 질문에 대한
답을 직접 도출하거나 결정짓는 주요 수단이 된다. 반
면, 다른 근거 자료가 존재할 경우 AI 모델은 보조적 수
단으로 다른 정보와 함께 해당 판단을 지원하거나 보완
하는 역할을 수행한다.

3단계

모델 위험성(Model Risk) 평가

이 단계에서는 모델의 위험성을 평가한
다. 여기서 말하는 ‘모델 위험성’은 AI 모델에 내재된 기
술적 위험성을 의미하는 것이 아니라, AI 모델의 출력
결과가 잘못된 판단을 초래해 바람직하지 않은 결과로
이어질 위험성을 의미한다. 이러한 모델 위험성은 두 가
지 요소를 고려해 평가할 수 있다.

- **모델 영향력 위험성(Model Influence Risk)**
본 지침은 이 위험성을 “AI 모델로부터 도출된 근거
가 질문에 대한 판단을 내리는 데 사용되는 다른 근
거와 비교했을 때 어느 정도의 영향력을 미치는지”
평가하는 것으로 정의한다. 즉 AI 모델이 의사결정
에 어느 정도의 영향을 미치는지를 평가하는 것이
다. (2단계에서 AI 모델이 보조적 수단에 그치는 경
우, 해당 AI 모델은 상대적으로 낮은 모델 영향력 위
험성을 가진 것으로 평가된다.)
- **결정 결과 위험성(Decision Consequence Risk)**
본 지침은 이 위험성을 “질문에 대한 잘못된 판단이
바람직하지 않은 결과로 이어질 가능성과 그 심각
성”으로 정의한다. 즉 AI 모델의 출력 결과가 부정확
할 경우 발생할 수 있는 잠재적인 영향 또는 위해 수
준을 평가하는 것이다.

결과적으로 위 두 가지 위험성 요소를 조합한 매트릭스
를 통해 모델의 위험성 수준(낮음, 중간, 높음)을 평가
할 수 있다.

위험성 평가를 위한 매트릭스

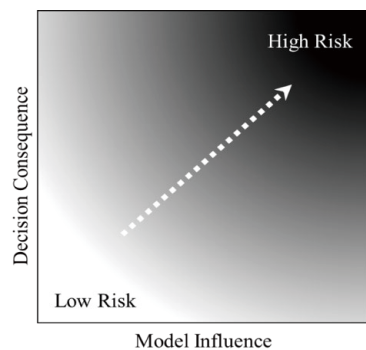


Figure 1. Model risk matrix. The model risk moves from low to high as decision consequence or model influence increases. The ratings for decision consequence and model influence are independently determined.

자료 FDA의 ‘의약품 및 생물학적 제품 관련 규제 의사결정을 지원
하기 위한 AI 사용에 대한 고려사항’ 지침(초안)

통상적으로 사람이 개입하지 않고 AI 모델이 단독으로
질문에 대한 최종 결정을 내리는 경우, FDA는 해당 AI
모델 사용을 높은 위험성을 가진 것으로 간주한다. 반
대로, 최종 결정에 이르기까지 AI 모델 외에도 사람의
검토와 평가가 요구되는 경우, 해당 AI 모델 사용은 낮

은 위험성을 가진 것으로 평가할 가능성이 높다.
특히 의약품의 안전성 또는 품질 등에 영향을 미칠 수
있는 고위험성 AI 모델의 경우, 다음과 같은 항목에 대
한 상세한 문서화가 요구될 수 있다.

- 모델 설계구조(model architecture)
- 데이터 출처
- 학습 및 튜닝 방법
- 검증 및 테스트 절차
- 신뢰구간(confidence intervals)을 포함한 성능 지표

반면, 저위험성 모델의 경우에는 상대적으로 간소한 수
준의 문서 제출만으로도 충분할 수 있다. 이와 같은 위
험성 기반 접근 방식은 중대한 규제 결정을 신뢰할 수
있는 근거에 기반하여 이루어지도록 하는 한편, 낮은
위험성 시나리오에서의 불필요한 규제 부담은 최소화
할 수 있다.

4단계

신뢰성 평가 계획(Credibility Assessment Plan) 수립

제약사는 AI 모델의 신뢰성을 어떻게 평가할 것인지에
대한 계획을 수립해야 한다. 해당 계획은 Context of
Use 및 3단계에서 평가된 모델의 위험성 수준을 고려
하여 수립해야 하며, 다음의 두 가지 핵심 요소를 포함
해야 한다.

- 모델 자체 및 모델 개발 과정에 대한 설명
- 모델 평가 과정에 대한 설명

모델 자체 및 모델 개발 과정에 대한 설명은 다음 사항
들을 포함해야 한다.

- 사용된 각 모델과 그 선택에 대한 근거, 입력 및 출력
에 대한 설명, 모델 설계구조, 특징(features), 특징
선택 과정, 매개변수(parameters) 등의 세부 정보
- 모델 개발에 사용된 개발 데이터(development
data): 학습 데이터(모델을 구축하기 위한 절차 및
알고리즘에 활용)와 튜닝 데이터(훈련된 소수 모델
을 평가하는 데 활용)로 구분됨
- 모델의 학습 방식, 학습 기법, 성능 지표, 정규화 기
법, AI 모델 보정(calibration) 방법, 품질 보증 및 품
질관리 절차 등

모델 평가 과정에 대한 설명은 다음 사항들을 포함해야
한다.

- 테스트 데이터가 AI 모델 평가를 위하여 어떻게 수
집, 처리, 주식 저장, 통제, 사용됐거나 또는 될 예정
인지에 대한 설명
- 데이터 독립성 확보 방식에 대한 설명
- 테스트 데이터가 Context of Use에 적합한지에 대
한 평가
- 모델의 예측 결과와 실제 관측 데이터 간의 일치도
- 선택한 모델 평가 방법의 타당성 근거
- 모델 평가에 사용된 성능 지표
- 편향 가능성을 포함한 접근 방식의 한계
- 품질 보증 및 품질 관리 절차 등

해당 계획을 FDA에 제출해야 하는지에 대한 여부, 제출
시기, 방식은 AI 모델의 위험성과 Context of Use뿐 아
니라 FDA와 어떻게 협의하는지에 따라 달라진다. 예컨
대, 해당 계획을 공식 회의 자료(formal meeting
package)의 일부로 포함하도록 요청받을 수도 있고, 내
부적으로 보관하다가 FDA의 요청이 있을 경우에만 제
출하도록 요구받을 수도 있다. 본 지침은 제출 시기나 구
체적 제출 요건에 대해서는 별도로 명시하고 있지 않다.

5단계

계획 실행

수립된 평가 계획은 데이터의 무결성, 검
증 방법, 문서화 등에 세심한 주의를 기울이며 실행돼
야 한다. 본 지침은 계획 실행 이전, 가급적이면 이른 시
점에 FDA와 협의할 것을 권장하며, 이를 통해 FDA와
제약사 간의 기대치를 유연하게 조율할 수 있도록 한다.

6단계

신뢰성 평가 보고서(Credibility Assessment Report) 작성

계획의 실행 결과는 신뢰성 평가 보고서로 정리돼야 하
며, 기존 계획과 달라진 사항(deviations)이 있을 경우
변경사항과 그에 대한 정당화도 함께 포함해야 한다. 이
보고서는 해당 Context of Use에 대한 AI 모델의 신뢰
성을 입증하는 근거로 작용한다.

7단계

모델 적정성(Model Adequacy) 평가

6단계에서 작성된 보고서를 바탕으로,
제약사는 해당 모델이 Context of Use에 대해 충분한

신뢰성을 갖추었는지를 판단해야 한다. 만약 제약사가 해당 모델이 충분한 신뢰성을 확보하지 못했다고 판단하는 경우, 아래와 같은 대응방법을 고려할 수 있다.

- 질문에 답하기 위한 다른 종류의 추가적인 근거 자료를 도입해 모델 영향력 위험성을 낮추는 방안
- 신뢰성 평가 활동의 엄격성을 강화하거나, 추가적인 개발 데이터를 통해 모델 출력을 보완하는 방안
- 위험성을 완화하기 위한 적절한 통제 방안 수립
- 모델링 접근 방식 자체를 변경하는 방안
- 모델의 Context of Use를 수정하는 방안



AI 모델의 생애주기(Life Cycle) 유지관리 필요성

본 지침은 AI 모델의 신뢰성은 유동적이며, 특히 의약품 생애주기 전반에 걸쳐 사용되는 경우, 지속적으로 유

지, 관리돼야 한다는 점을 강조한다. AI 모델은 데이터 드리프트와 같은 데이터 입력값의 변화, 배포 환경의 변화, 또는 제조 공정의 변경 등에 따라 성능 저하에 취약할 수 있다. 이에 따라 제약사는 AI 모델이 해당 Context of Use에서 의도한 바의 기능을 지속적으로 수행할 수 있도록 지속적인 모니터링과 재평가 계획을 수립해야 한다.



제출 요건

본 지침은 AI를 활용하는 제약사가 FDA와의 협의를 통해 정해지는 AI 모델 관련 정보를 기존 요구되던 자료에 더해 추가로 제출해야 한다는 점을 강조한다. 이는 해당 정보가 기존 자료를 대체하는 것이 아니라, 별도의 추가 제출 자료임을 의미한다. 이는 AI를 활용하는 이해관계자들이 의약품 자체를 뒷받침하는 정보뿐 아니라 신약개발 또는 규제 관련 제출 과정에 기여한 AI 모델의 신뢰성과 성능을 입증하는 데 필요한 정보까지 포괄적으로 준비하는 것을 요구한다.

앞서 언급한 바와 같이 제출 여부, 시기, 방식 및 제출 자료의 구체적 범위 및 내용은 AI 모델의 위험성과 Context of use를 고려한 FDA와의 사전 협의 결과에 따라 달라질 수 있다.



전략적 시사점

① 개발 초기부터 FDA와 선제적으로 협의하라

본 지침 전반에 걸쳐 FDA는 선제적인 초기 협의의 중요성을 강조하고 있다. AI의 활용 사례는 복잡하고 개별 Context of Use에 따라 다르기 때문에, 지침은 일률적 제출 요건을 강제하지 않는다. 따라서 제약사가 초기에 FDA와 협의할 경우 FDA가 해당 AI 모델의 Context of Use 및 평가된 위험성 수준에 맞춰 제출 요건 및 피드백을 보다 유연하고 구체적으로 조정할 수 있다. 이는 보다 효율적인 규제 이행을 가능하게 하고, 불필요한 정보 공개의 부담도 줄일 수 있다.

제약사는 INTERACT 미팅, 사전 IND(Pre-IND) 미팅 등 다양한 절차를 통해 FDA와 협의할 수 있다. 본 지침은 AI 모델의 적용 분야에 따라 활용할 수 있는 여러 맞춤형 협의 경로를 제시하고 있다.

② 문서화는 신약개발 과정에 내재화하라

FDA가 요구하는 정보에 효과적으로 대비하기 위해서는 문서화 전략이 개발 과정 전반에 내재되어야 한다. 특히, AI 모델이 고위험성으로 분류될 가능성이 있는 경우에는 문서화의 범위와 정밀도에 대한 요구 수준이 높아지기 때문에, 개발 초기부터 체계적으로 준비하는 것이 중요하다. 여기서의 '문서화'란, 본 지침의 7단계 절차에 따라 각 단계별로 요구되는 분석, 자료, 평가 결과 등을 규제 목적에 맞게 정리하고 관리하는 전반을 의미한다. 따라서 문서화는 신약개발 완료 후에 사후적으로 작성하는 정리 작업이 아니라, 신약개발 생애주기 전반에 걸쳐 내재화돼야 하는 핵심 요소다.

③ FDA의 정보공개 요구 수준에 맞춘 지식재산권(IP) 전략을 수립하라

제약사는 본 지침이 요구하는 투명성 수준이 자사의 AI 모델 관련 기술 보호 방식에 영향을 미칠 수 있다는 점을 염두에 두어야 한다. FDA는 AI 모델의 설계구조, 데이터, 학습 방식 등에 대해 상세한 공개를 요구할 수 있으므로 AI 모델의 독점적 요소들이 외부에 노출될 수 있다. 이 경우 해당 정보를 미국법상 영업비밀(trade secret)로 보호하기 어려워질 수 있기 때문에 제약사는 기존 영업비밀 전략에 일부 의존하고 있었다면 해당 전략이 여전히 유효한지, 아니면 공개를 전제로 한 특허권 보호가 더 효과적인 수단인지를 포함해 지식재산권(IP) 전략 전반을 재검토할 필요가 있다. 이는 궁극적으로 FDA가 어느 정도 수준의 정보 공개를 요구하는지에 따라 달라질 것이다.

또한 AI 모델 거버넌스, 위험성 관리 체계, 검증 방법론 등은 단순 규제 대응 차원을 넘어, 특허 보호가 가능한 독립적 기술 자산으로 평가될 수 있다. 이는 곧 제약사에게 새로운 사업기회를 의미한다.



앞을 내다보며: 산업과 발맞춰 진화하는 규제환경

지난 1월에 발표된 본 지침은 FDA가 단순히 산업의 흐름을 따라가는 것을 넘어, AI가 신약개발에 미치는 영향을 선제적으로 규정하고 이끌고자 하는 의지를 보여주는 일련의 최근 조치 중 하나다. 지난 2월, FDA는 CBER(Center for Biologics Evaluation and Research), CDER(Center for Drug Evaluation and Research),

CDRH(Center for Devices and Radiological Health), OCP(Office of Combination Products) 등 FDA 내 주요 센터 간 협력에 관한 내부 정책 문서 '인공지능과 의약품의료 제품: CBER, CDER, CDRH 그리고 OCP의 협력 체계(Artificial Intelligence and Medical Products: How CBER, CDER, CDRH, and OCP Are Working Together)'를 개정하며 AI 모델 성능 모니터링, 생애주기 관리, 편향 완화 등에 대해 보다 구체적이고 강한 조의 내용을 추가했다. 더불어 지난 4월에는 전통적인 동물실험 요건을 단계적으로 축소 및 폐지하고 AI 기반 시뮬레이션 및 인체 유사 장기 모델 등 대체 시험방법 도입을 장려하겠다고 발표했다.

이와 함께 FDA는 내부적으로도 AI 도입을 적극 추진하고 있다. FDA는 2024년 8월에 CDER의 AI 위원회를 공식 출범한 것에 이어, 지난 5월에는 첫 최고 AI 책임자(Chief AI Officer)를 임명했다. 이와 비슷한 시기에 FDA가 생성형 AI 기반의 과학 심사 도구 개발을 위해 오픈에이아이(OpenAI)와 예비 협의를 진행 중이라는 보도가 있었으며, 해당 시스템은 'cderGPT'라고 불리는 것으로 전해진다.

같은 5월, 마틴 마커리 FDA 국장은 FDA가 생성형 AI를 활용한 첫 과학적 심사 작업을 성공적으로 운영했으며, 2025년 6월 30일까지 전 부서로의 도입을 목표로 한다고 밝혔다. 6월 초에는 '엘사(Elsa)'라는 이름의 생성형 AI가 공식 도입되었으며, 이는 심사관과 조사관을 포함한 FDA 직원들의 업무를 지원하고, FDA 내부 운영을 최적화하며 향후 AI 통합의 기반을 마련하는 역할을 하도록 설계되었다고 발표되었다.

이러한 일련의 조치는 FDA가 단순히 업계의 빠른 AI 기술 채택에 수동적으로 대응하고 규제체계를 마련하는 수준을 넘어, 자체적으로도 AI를 수용하고 활용하는 방향으로 조직을 재편하고 있다는 점을 분명히 드러낸다. 미국시장 진출을 계획하고 AI를 신약개발에 활용하고자 하는 국내 제약사는 반드시 미국 규제환경의 방향성과 기대 수준을 전략적으로 이해하고 선제적으로 대응해야 한다. FDA와의 초기 협의와 규제 대응 전략을 어떻게 설계하느냐에 따라, AI 기술 기반 혁신이 규제 장벽이 될 수도, 경쟁 우위가 될 수도 있다. ①

*본 지침은 FDA가 공공 의견 수렴을 목적으로 발행한 초안 형태의 문서다. FDA는 발표일로부터 90일간(2025년 4월 7일까지) 대중의 의견을 접수했다. 접수된 의견들은 FDA가 향후 최종 지침을 확정하는 과정에서 검토, 반영될 것으로 예상된다.

국내 제약사는 반드시 미국 규제환경의 방향성과 기대 수준을 전략적으로 이해하고 선제적으로 대응해야 한다.

