

Deloitte Insights

Aug 2024



생성형AI와 정부의 업무 혁신 2편 : 정부 규제 운영 업무 혁신 방법

Deloitte Center for Government Insights

Deloitte.

Download on the
App Store

GET IT ON
Google Play



'딜로이트 인사이트' 앱에서
경영·산업 트렌드를 만나보세요!

목차

· AI와 규제 운영 업무 혁신	03
· 규제 운영에서 AI의 잠재력을 실현하기 위한 세 가지 고려 사항	04
(1) 다양한 도구의 고유한 기능을 이해하기	04
(2) 직무에 적합한 복수의 도구 찾기	04
(3) AI와 인간의 판단을 통합하도록 비즈니스 프로세스 조정하기	05
· AI가 규제 운영을 혁신하는 방법	06
· AI를 활용해 공개 의견 조작에 대응하기	06
· 기관의 기억(institutional memory) 유지	08
· 실태조사의 효과성 향상	08
· AI 도입 프로그램 착수	09
(1) 기본적인 기술에 익숙해지기	09
(2) 사용자를 중심에 두기	09
(3) AI에 대한 신뢰 구축에 투자하기	09
· 더 좋고 빠른 규제 프로세스의 미래를 향해	10

AI로 정부 규제 운영 업무를 혁신하는 방법

강력한 인공지능(AI) 도구를 갖춘 정부 규제 기관은 기업과 시민 모두와 협력하고 상호 작용하는 방식을 바꿀 수 있다.

AI와 규제 운영 업무 혁신

수십 년 동안 경제는 디지털화되었다. 오늘날 항공기는 동체를 이루는 물리적 부품보다 소프트웨어 코드 갯수가 더 많다.¹ 완벽한 가상 환경에서 제작된 영화를 디지털 테이프에 저장할 수 있게 되었다.² 그리고 금융 거래가 너무 빨리 이루어지다보니 일부 증권 거래소는 긴 광섬유 케이블을 이용해 이러한 거래 속도를 늦춰야 할 지경이다.³

이렇게 여러 산업들이 놀라운 속도로 디지털화되고 있기는 하지만, 그래도 여전히 안전하게 운영하는 것이 중요하다. 항공기는 안전해야 하고, 영화의 저작권은 보호되어야 하며, 금융 거래의 안정이 유지되어야 한다. 그러나 정작 이러한 산업을 감독하는 규제 기관은 아직도 구식 도구를 사용해 효율성이 뒤쳐진다. 예를 들어, 딜로이트가 AI를 사용하여 2017년 미국 연방규정집(2017 CFR)의 21만7,000개가 넘는 섹션을 분석한 결과, 중복되거나 유사한 구절이 포함된 섹션이 거의 1만8,000개에 달한다는 것을 발견했다.⁴ 현대 규제 기관은 현대적인 도구를 사용해야 한다. AI가 중복 규정을 식별하는 데 도움이 된다면, 중복되거나 상충되는 규정을 제거하는 막대한 기회를 창출할 수 있다. 게다가 이 정도는 시작에 불과할 수 있다. 최근에 방대한 양의 텍스트를 처리할 수 있는 대규모 언어 모델(LLM)과 다른 형태의 생성형AI가 등장하면서 일관적이고도 맥락에도 잘 들어 맞는, 심지어 예술적인 결과물까지 만들어낼 수 있는 강력한 새로운 도구를 쓸 수 있게 됐다.⁵

이렇게 강력한 AI 도구를 갖춘 규제 기관은 기업 및 시민 모두와 협력하고 상호 작용하는 방식을 바꿀 수 있다.⁶ 예를 들어, 생성형AI는 시민 및 기업과의 상호 작용을 개선하고, 방대한 양의 이해관계자 의견을 분석 및 요약하고, 보고서 생성과 같은 행정 업무를 자동화하고, 소프트웨어 솔루션을 코딩하고, 심지어 맞춤형 솔루션을 제시할 수도 있다.

AI가 모든 문제에 해답을 제공하는 것은 아니지만, 알맞은 작업에 적용되어 적절한 인간의 판단과 결합될 경우 솔루션의 매우 중요한 부분을 담당할 수 있다.



규제 운영에서 AI의 잠재력을 실현하기 위한 세 가지 고려 사항

생성형AI의 잠재력을 활용하고자 하는 규제 기관은 다음과 같은 단계적인 조치를 고려해야 한다. 1)다양한 AI 도구의 고유한 기능을 이해하기 2)고유한 직무에 적합한 복수의 도구 채택하기 3)AI와 인간의 판단을 통합할 수 있도록 비즈니스 프로세스 조정하기.

이러한 원칙을 수용하면 규제 기관은 규제 업무 처리를 빠르게 하고, 직원의 부담을 줄이기 위해 내부 프로세스를 자동화하고, 데이터 기반 의사 결정을 통해 규정 준수율을 개선하는 등 감독하는 산업과 보조를 맞추는 데 도움이 될 수 있다.

(1) 다양한 AI 도구의 고유한 기능을 이해하기

서로 다른 AI 도구는 차별적인 방식으로 작동한다. 예를 들어 생성형AI는 다른 형태의 머신러닝(ML)이 할 수 없는 창의적인 작업을 할 수 있지만, 어느 정도 정확성이 떨어지는 단점이 있다. 예를 들어, 생성형AI가 그럴 듯하지만 잘못된 사실로 응답하는 환각(hallucination) 문제에 대해 익히 들었을 것이다. 반면 보다 전통적인 형태의 머신러닝은 데이터에서 계층적 관계를 식별할 수는 있지만, 왜 이러한 특정 결론에 도달했는지를 쉽게 보여주지 못한다.

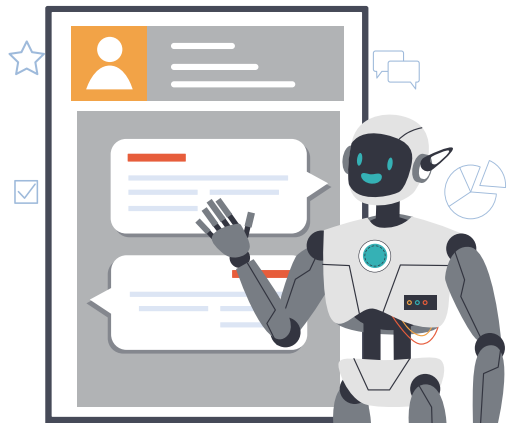
이러한 강점과 약점의 상호 작용 때문에 어떤 AI 도구가 다른 도구에 비해 차별적인 직무에 보다 더 적합할 수 있다. 청문회의 진술을 기록하거나 보고서를 생성하거나 통화에서 사용자 감정을 분석하는 것과 같이 반복적으로 콘텐츠를 생성해야 하는 작업은 생성형AI에 맡길 수 있다. 중복된 법률을 식별하거나 사기를 감지하는 것과 같이 높은 수준의 정확도로 대량의 데이터에서 패턴을 찾는 작업은 기존 머신러닝 기술에 할당하는 것이 좋다. 마지막으로, 높은 가변성을 가지거나 사회적인 구성 요소를 처리할 때 인간의 판단이 가진 강점은 AI가 이러한 인간의 전문성을 강화하는데 기여해야 한다는 것을 의미한다. AI 도구를 적용하면 인간이 더 어려운 작업을 보다 효율적으로 수행할 수 있다.

덴마크 기업청(DBA)은 이미 이와 유사한 방식의 시스템을 도입했다. 우선 AI를 사용하여 연간 23만 건 이상의 재무제표를 분석해 사기 가능성과 불일치 사항을 파악하는, 계산 업무량이 매우 막중한 작업을 수행한다.⁷ 하지만 일단 문제점이 발견되면, 여기서 사기 행위가 이루어졌는지 여부를 결정할 때는 인간의 판단이 제공하는 것과 같은 내용물에 대한 인식이 필요하다.

(2) 직무에 적합한 복수의 도구 채택하기

다양한 AI 모델은 서로 다른 작업을 하기 때문에, 가장 정교하고 값비싼 도구를 채택해도 최상의 결과를 얻지 못할 수 있다. 최상의 결과는 각각 고유한 강점을 발휘하는 여러 개의 작은 도구를 활용할 때 얻을 가능성이 높다. AI 증강형 세무 운영을 살펴보자. AI 디지털 비서를 이용해 5분 만에 세금 신고서를 준비할 수 있다.⁸ 자연어처리(NLP)가 가능한 챗봇은 납세자에게 세무 관련 질문을 하고 로봇 프로세 자동화(RPA)의 도움으로 세무 보고서 항목을 자동으로 기입해줄 수 있다. 마지막으로 납세자가 세무 보고서를 검토하고 데이터를 확인한 다음 보고서를 제출하면 된다. 나아가 챗봇에게는 간단한 세무 보고서 작성을 맡기고, 더 복잡한 보고서는 인간이 담당할 수 있다.

단일한 생성형AI 도구의 경우 너무 창의적인 세무 자문을 하지 않도록 매우 엄격한 매개변수가 필요하다. 각각의 강점을 살린 다양한 AI 도구를 결합하면 단일한 거대 도구보다 적은 자원으로 원하는 결과를 생성할 수 있다.



(3) AI와 인간의 판단을 통합하도록 비즈니스 프로세스 조정하기

종이서류 업무 시대에 구축한 규제 프로세스에 새로운 기술을 추가하면 역효과가 날 수 있다. 규제 기관은 AI가 제공할 수 있는 것이 무엇인지 이해하고 규제 프로세스를 재고해야 한다. 딜로이트의 설문 조사에 의하면, AI 프로젝트가 중대한 업무 절차 변경을 수반할 경우 그 성공 가능성이 36% 더 높은 것으로 나타났다.⁹

효과적인 프로세스는 먼저 규제 결과를 결정한 다음 관련 직무를 수행하는 데 가장 적합한 기술을 식별하는 방식으로 구축할 수 있다. 미국 증권거래위원회(SEC)는 투자자문사의 규제 신고 문서를 분석하기 위해 NLP와 머신러닝을 통합하는 도구에 투자했다. SEC는 이것으로 기존 프로세스를 더 빠르게 만드는 것이 아니라, 신고서에서 잠재적인 위반을 적발할 수 있도록 특정 패턴을 찾는 것과 같은 새로운 방식을 시도했다. 과거 데이터를 시뮬레이션하는 백테스팅(backtesting) 분석 결과, 알고리즘은 신고서에 대한 무작위 검사보다 5배나 더 뛰어나게 집행기관에 회부할 만한 언어를 식별해 냈다.¹⁰

이러한 모니터링 도구는 규정 위반 가능성 또는 그러한 영향이 가장 높은 곳을 식별하여 규정 준수 활동에 대한 위험 기반 표적화(risk-based targeting)를 가능하게 해준다. 다시 말해, 이러한 도구를 사용하면 규제 기관이 차별화된 방식으로 업무를 수행할 수 있다. 예를 들자면, 위반이 발생했을 때를 기다렸다가 벌금이나 시정 조치를 내리는 대신, 신고서류를 선제적으로 검사할 수 있는 것이다.



시가 규제 운영을 혁신하는 방법

앞서 세 가지 고려 사항만 봐도 시가 규제 운영에 엄청난 영향을 미칠 수 있다는 것은 분명하다. AI는 기존 업무의 효율성을 개선하는 데 도움이 될 수 있다. 덴마크 기업청이 AI를 사용하여 사기를 식별하는 것이 그것이다. 그러나 규제 기관은 SEC가 NLP와 AI를 사용하여 투자자문사의 위반 행위를 감지하는 것처럼 근본적으로 새로운 방식으로 원하는 결과를 달성할 수도 있다. 또한 AI는 공개 협의, 규제 개발 및 규제 시행 등을 포괄하는 규제 주기(regulatory life cycle) 전반에 걸쳐 효율성과 효과성 모두를 개선하는 데 도움이 될 수 있다(그림 1).

그림 1. AI는 규제 주기 전반에 걸쳐 효율성과 효과성 모두를 개선할 수 있다

아래 사례들은 시가 규제 주기를 혁신하도록 도울 수 있다는 것을 보여준다

● 시가 효과성 개선 ● 시가 효율성 개선

● 여론 예측 ● 이해관계자 참여 챗봇	● 규제 초안 작성 ● 적응형 규제 도구 코딩	● 조사 대상 목록 자동 생성 ● 승인 및 거절 응답 생성	● 정책 효과 평가 및 시뮬레이션 ● 상충 규정 식별 및 수정
● 시가 생성한 공개 의견 선별 ● 입력물 분석 및 요약	● 중첩 및 중복 규제 식별 ● 정책 및 규제 시나리오 분석	● 예측 및 위험기반 집행 ● 사기 탐지	● 규제 성과 현황판 ● 규제준수율 추적
협상 및 참여	규제 및 법률 개발	규제 시행	규제 검토
규제 주기			

출처: 딜로이트 분석

AI를 활용해 공개 의견 조작에 대응하기

의견 청취 기간은 대다수 규제 의사 결정에서 중요한 요소다. 그러나 의견 청취의 중요성과 개방성은 조작에 취약하게 만든다. 미국 연방통신위원회(FCC)가 2017년에 ‘망 중립성’에 대한 새로운 규칙을 제안했을 때, 온라인 의견 청취 시스템은 2,200만 건 이상의 댓글을 받았다. 연구자들은 NLP 및 클러스터링 알고리즘과 같은 도구를 사용하여 그 댓글들 중에서 100만 건 이상, 혹은 그보다 훨씬 더 많은 댓글이 봇(bot)이 쓴 것이라는 사실을 밝혀냈다.¹¹ FCC는 댓글 내에서 반복되는 문구, 단어 사용 및 문장 구조를 검색하기 위해 계약업체를 고용했다. 이들은 동일한 문장과 단락 내에서 단어를 동의어로 바꾼 ‘고유한’(unique) 메시지를 최대 130만 건 발견했다.¹² 특정 옹호 단체가 작성한 메시지를 실제로 사람들이 복사하여 붙여넣기하는 이메일 캠페인과 달리, 이러한 메시지는 각각 아마도 봇으로 추정되는 주체가 약간씩 변경하여 걸러내기 장치를 우회하고 ‘고유한’ 것처럼 보이도록 했다. 한 분석에 따르면 스팸 메시지나 대규모 이메일 캠페인의 일부가 아니라 자연스럽게 작성된 댓글은 80만 건에 불과했다.¹³

봇이 생성한 압박 캠페인에 맞서기 위해 조직의 리더는 다양한 도구를 사용하여 댓글을 분류할 수 있다. 2,200만 개의 댓글을 모두 읽는 것은 인간에게는 불가능하지만 AI는 댓글을 그룹으로 묶어 수백만 개의 댓글을 25개 정도의 대량 메시지로 축소할 수 있다. 상위 25개

댓글은 망중립성 폐지에 찬성하는 댓글의 98%를 차지했다.¹⁴ '감정 분석'(sentiment analysis)을 통해서는 댓글을 긍정 및 부정으로 분류하여 대략적인 찬반 조사 결과를 제공할 수 있다.

또한 AI를 활용해 또한 규제 기관이 공개 의견 청취 기간의 프로세스를 재구성할 수 있다. 스팸 감지는 봇 공격과 AI가 생성한 의견을 걸러내는 데 도움이 될 수 있고, NLP와 감정 분석은 적법한 의견을 요약할 수 있다. 생성형AI는 이들의 주장을 구분하고, 심지어는 관련된 지지층 대해 추론할 수도 있다. 마지막으로 생성형AI는 정책 입안자가 가장 중요한 문제에 집중할 수 있도록 여러 이해 관계자의 토론 결과를 요약할 수 있다.

적절한 안전장치를 활용하면 AI는 댓글의 이러한 공개 의견 내용뿐만 아니라 메타데이터도 평가할 수 있다. 운영 체제, IP 주소, 제출 시간, 브라우저와 같은 요소들은 모두 스팸을 식별하는 데 도움이 된다. 한 데이터 분석 회사는 "\n 문자열"이라고 하는 보이지 않는 줄 바꿈을 기반으로 한 의심되는 FCC 봇 댓글을 발견했다.¹⁵

AI는 시스템이 개인 정보를 보호하기 위해 노력하는 한에서 분석에 이러한 요소들을 포함할 수 있다. 데이터 사용 방법에 대한 투명성과 제한을 요구하는 딜로이트의 '신뢰할 수 있는 AI' 프레임워크(아래 박스의 "신뢰할 수 있는 AI의 6가지 차원" 참조)는 AI가 개인 정보를 침해하지 않고 관련 데이터를 평가하는 데 도움이 될 수 있다.¹⁶

신뢰할 수 있는 AI의 6대 요소

딜로이트의 신뢰할 수 있는 AI 프레임워크는 AI에 대한 신뢰를 구축하는 데 도움이 되는 6가지 주요 요소를 제시한다. 이 프레임워크는 규제 기관이 AI와 관련된 잠재적 위험을 식별하고 완화하는 데 도움이 되도록 설계되었다.

- ✔ ① **공정성 및 공평성:** AI는 공정하고 일관된 프로세스를 따르도록 설계되고 훈련되어야 하며, 공정한 결정을 내려야 한다. 허가나 면허를 승인하거나 거절하는 것과 같은 규제 결정을 포함하여 차별적 편견을 줄이기 위한 내부 및 외부 검사를 포함해야 한다.
- ✔ ② **투명성 및 설명 가능성:** 사용자는 특히 의사 결정에서 기술이 어떻게 활용되는지 이해해야 한다. 규제 기관은 해당 알고리즘의 영향을 받는 사업체에 투명하고 설명할 수 있는 알고리즘 생성의 중요성을 강조해야 한다.
- ✔ ③ **책무 및 책임성:** AI 시스템은 시스템의 산출물 또는 의사 결정에 대해 누구에게 책무가 있으며 책임을 지는지 설명하는 정책을 가져야 한다. 이는 생성형AI가 법률 초안 작성 및 정책 초안에 대한 이해 관계자의 의견을 요약하는 것과 같은 규제 응용 프로그램에서 사용됨에 따라 점점 더 중요해질 것이다.
- ✔ ④ **안전성과 보안성:** AI를 신뢰할 수 있으려면 시스템을 조작하고 디지털 피해를 입힐 수 있는 사이버보안 위험으로부터 보호되어야 한다.
- ✔ ⑤ **개인정보보호 존중:** AI 시스템에서 생성한 정교한 통찰력은 종종 세부적이고 개인적인 데이터에서 나오기 때문에 개인정보 보호가 매우 중요하다. 신뢰할 수 있는 AI는 데이터 규정을 준수해야 하며 명시되고 합의된 목적으로만 데이터를 사용해야 한다.
- ✔ ⑥ **강건성 및 신뢰성:** AI는 적어도 자신이 증강하거나 대체하는 기존 시스템, 프로세스 또는 사람만큼 강건하고도 신뢰할 수 있어야 한다. 특히 AI가 결과물을 확장될 때 일관되고 신뢰할 수 있는 산출물을 생성해야 한다.

기관의 기억(institutional memory) 유지

규제 기관은 종종 인간의 기억에 의존한다. 규제 당국은 규제 전략을 시험하고, 허점을 메우고, 규정을 작성한다. 규제 당국자는 자신이 담당하는 기술, 시설 및 산업의 전문가가 된다. 그들은 정부의 절차를 탐색하는 동안 꾸준하고 예측 가능한 서비스를 제공하는 방법을 배운다. 오랫동안 근무한 규제 전문가가 소속 기관을 떠날 경우, 이러한 방대한 지식의 손실이 발생할 수 있다. 또한 이직을 통해 새롭게 들어온 규제 및 정책 담당 직원이 업무를 시작하기 어려워 규칙 제정이나 규제 시행이 지연될 수도 있다.

이러한 문제를 완화하는 데 AI가 도움이 될 수 있다. 먼저, 직원이 퇴사할 때 퇴사 면접에서 '암묵 지식'을 포착할 수 있다. 그런 다음 생성형AI가 이러한 면접 결과를 활용하여 신입 직원의 특정 직무에 맞게 조정된 온보딩 및 교육 문서를 만들 수 있다. 생성형AI는 기존 법률 및 규정, 과거 시행 결정에 대해 훈련을 받을 수 있으므로, 기존 직원과 신입 직원은 채팅봇에 질의를 통해 복잡한 정책 문서를 더 잘 이해하고 규제 연구 프로세스를 가속화할 수 있다.

실태조사의 효과성 향상

AI는 실태조사 우선순위 지정, 조사 중 서류 작업 감소, 보고서 초안 작성, 과거 패턴에 대한 통찰력 발견 등 조사 과정 전반에서 업무를 개선하기 위한 다양한 선택지를 제공한다.

조사 목표를 지정하는 도구는 잘 확립되어 있다. 일부 보건부서는 AI를 사용하여 어떤 식당이 규정을 위반할 가능성이 가장 높은지 사전에 파악한다.¹⁷ 해당 프로그램은 먼저 쓰레기에 대한 311건의 불만에서 식중독에 대한 트위터(Twitter) 게시물에 이르기까지 모든 것을 사용하여 위험 요소를 평가한다.¹⁸ 이러한 위험 평가는 위반 가능성이 더 높은 조사 대상을 지정하는 데 도움이 되며, 무작위 검사를 실시한 대조군보다 위반 적발 건수가 3분의 1 이상 더 많았다.¹⁹ 연구자들은 이러한 소프트웨어가 연간 9,000건의 식중독 사고와 500건 이상의 병원 내원을 줄일 수 있다는 추정 결과를 제출했다.²⁰

마찬가지로 뉴욕시의 건물 검사관은 AI 도구를 사용하여 어떤 요인이 '구조물 화재' 발생과 관련이 있는지 확인하고, 모든 건물에 위험 점수를 할당했다. 이러한 도구는 잠재적으로 위험한 식당을 대상으로 하는 보건 검사관 도구와 매우 유사한 방식으로 건물 조사를 특정 대상으로 집중하는 데 도움이 된다. 이 도구는 학교와 같이 의무적인 연간 검사를 받는 위험이 낮은 건물들을 통합하는 것을 포함하여 우선순위에 따라 건물을 자동으로 배열한다. 뉴욕시 소방당국은 처음에는 약 6가지 주요 위험 요소를 고려했지만, 최근 반복 작업을 통해 17개 시 기관의 데이터 스트림(data stream)에서 7,500개 이상의 데이터 포인트(data point)를 추적하고 있다.²¹

AI는 각각의 새로운 화재에서 학습하여 고려해야 할 새로운 요인을 식별하고 실시간으로 위험을 재평가한다. 이를 통해 소방당국 리더가 계획을 수립하는 데 활용할 수 있는 화재 위험의 실시간 현황판을 제공한다. 유사한 알고리즘이 직장 위험에서 환경 유출에 이르기까지 거의 모든 규제 영역에서 고위험 상황을 예측할 수 있다.

일단 조사가 시작되면 즉시 생성형AI를 적용할 수 있다. 제조 공장 시설에 대한 신규 인가 과정에서 대기 질 검사관이 규정이 제대로 적용되는지 이해하는 데 며칠이 걸릴 수 있다. 예를 들어, 해당 시설이 밤에 특정 규모보다 큰 가스 보일러를 가동하는 것에 대한 조항이 적용되는가? 생성형AI는 시설 인가를 검사할 핵심 사항으로 요약하고, 관련 규정의 섹션에 대해 강조점을 부여할 수 있다.

검사관은 물론 모든 내용을 읽어봐야 하지만 적어도 이러한 요점 정리부터 시작할 수 있다. 여기서도 '기관의 지식'이 필요하다. 특정 시설의 과거 위반 사항을 기억하는 챗봇은 어디를 살펴봐야 할지 제안할 수 있고, 과거 검사관으로부터 얻은 정보는 새로운 검사관에 게 모든 화학 물질을 보관하는 문 뒤를 확인해야 한다는 점을 환기할 수 있다. AI는 또한 모든 시설에 대한 브리핑 양식을 만드는 데도 도움이 될 수 있다.

조사 우선순위를 정하는 것 외에도 생성형AI는 위반 사항 문서화, 통지 발송, 조사 보고서 초안 작성에도 도움이 될 수 있다. 서류 작업에 드는 시간 절약은 덤이다. 2017년에 실시한 미국 연방 정부 인력 분석에 따르면, AI를 통해 수동 작업을 자동화하면 규제 담당 직원의

시간을 무려 수천만 시간 절약할 수 있을 것으로 추산했다. 예를 들어, 규정 준수 및 시행 작업과 관련된 활동만 해도 연간 최대 6,000만 시간을 절약할 수 있다. 검사관의 시간도 연간 2,600만 시간 절약될 수 있다.²²

이러한 분석이 이루어진 이후 6년 동안 AI의 역할은 크게 개선되었고, 생성형AI의 등장으로 인해 시간 절감 가능성은 더욱 높아졌을 것이다. 이제 AI가 훨씬 더 많은 낮은 가치의 작업을 수행할 수 있다면, 검사관은 더 높은 가치의 인지적 작업을 수행할 수 있다.

AI의 증강 기능을 통해 검사관은 본래의 업무, 즉 고위험 작업을 검사하는 데 더 많은 시간을 할애할 수 있다.

AI 도입 프로그램 착수

보완적인 AI 도구가 인간의 판단과 신중하게 결합된다면, 생성형AI는 규제의 개발 및 시행에 혁신적인 가능성을 열어줄 수 있다. 이러한 혁신으로 가는 길에는 몇 가지 중요한 해결과제가 있지만, 규제 기관이 AI 여정을 시작할 때 다음 세 가지 고려사항이 성공 가능성을 높이는 데 도움이 될 수 있다.

(1) 기본적인 기술에 익숙해지기

다양한 유형의 머신러닝과 생성형AI의 기본 기능을 이해하면, 조직의 리더가 어떤 기술이 어떤 작업에 적합한지 결정하는 데 도움이 될 수 있다. 기존 머신러닝의 기초 설정이 지닌 상대적인 투명성은 생성형AI 도구보다 사건 보고서에서 패턴을 찾는 데 좀더 나을 수 있다. AI 모델이 작동하는 방식에 대한 자세한 지식은 개인정보보호 및 보안 문제를 완화하는 데에도 도움이 될 수 있다. 예를 들어, LLM이 검토할 문서의 범위를 제한하기 위해 '지식 임베딩'(knowledge embedding)을 사용하면 정확도를 높이고 쿼리를 공개 도메인에서 제외하는 데 도움이 될 수 있다. AI의 기능을 알면 무엇이 가능한지 상상할 수 있다.

(2) 사용자를 중심에 두기

재구성된 규제 프로세스는 잘 설계된 모든 기술 프로젝트처럼 작동해야 한다. 즉 사용자를 중심에 두는 것이다. AI가 증강하는 프로세스에 따라, 사용자는 시민, 기업 또는 규제 기관 자체일 수 있다. 사용자를 중심으로 하면 좋은 디자인을 만들어 내는 데 도움이 될 수 있다. 여기에는 인간 중심 디자인(HCD)의 기본 원칙이 적용된다. 먼저 사용자가 원하는 결과에 초점을 맞추고, 현재 그 진행을 방해하는 요소가 무엇인지를 파악한다. 다음으로 기술을 도입하여 이러한 해결 과제를 줄이는 데 도움이 되는 좀더 정교한 솔루션을 개발한다. 마지막으로 임무를 수행하는 데 상관관계가 있는 지표를 찾아 지속적으로 이를 측정한다.

(3) AI에 대한 신뢰 구축에 투자하기

경계가 진화하면 규제 기관도 이와 함께 진화해야 하지만, 종종 규제는 느리게 회전하는 선박과 같다. 빠르게 움직이는 기술의 이점과 신뢰 및 안정성 사이의 균형을 맞추는 것이 규제 기관에게 어려운 과제가 될 수 있다.

규제 당국에게는 신뢰가 중요하다. AI를 신뢰하는 것과 관련하여, 딜로이트의 신뢰할 수 있는 AI 프레임워크의 핵심 원칙은 투명성이다. 신뢰할 수 있는 AI를 개발하려면 두 가지 측면에서 투명성이 필요하다.

첫째, AI 모델 자체의 투명성이다. 입력이 들어오고 결과가 나오는데 중간에 무슨 일이 일어났는지 아무도 확인할 수 없는 '블랙박스' AI 모델의 불투명성은, 특히 그 결과에 당신의 생명이나 생계가 걸려 있는 경우, 큰 우려의 대상일 수밖에 없다. 이해 관계자가 이해할 수 있는 감독 메커니즘이 있을 때 대중의 신뢰가 더 강해질 수 있다. 따라서 규제 기관이 AI를 사용하여 의사 결정을 강화하는 경우, 설명 가능

한 기능이 있는 알고리즘을 선택해야 한다. AI가 데이터에서 숨겨진 추세를 찾는 데 사용되는 경우, 모델을 학습하는 데 사용된 데이터의 기능은 학습 데이터 자체가 아니라면 공개되어야 한다. AI가 프로세스를 자동화하는 경우 결과는 일관되어야 한다. 정부의 AI는 민간 부문의 실험적 모델보다 더 높은 신뢰 기준을 충족해야 한다.

투명성의 두 번째 측면은 AI 모델의 목적에 관한 것이다. AI 프로그램은 공정하고, 투명하고, 책임감 있고, 일관되고, 안전해야 하며, 개인 정보를 보호해야 한다. 조직의 리더는 처음부터 모든 올바른 거버넌스 보호 장치를 구축할 뿐만 아니라 AI 모델을 언제, 왜 사용하는지 명확히 해야 한다. 이러한 명확성은 AI가 개인 정보를 침해하거나 근로자를 대체하는 데 사용되지 않고 대중에게 실질적인 혜택을 제공하도록 설계되었다는 것을 대중과 근로자에게 보여주는 데 도움이 될 수 있다. 딜로이트의 연구에 따르면 혜택에 대해 의사소통하는 것이 정부가 신뢰를 얻는 가장 확실한 방법 중 하나이다.²³

더 좋고 빠른 규제 프로세스의 미래를 향해

AI는 큰 변화를 가능하게 한다. 머신러닝을 통합하여 더 나은 프로세스를 만드는 리엔지니어링(reengineering) 전망이 부담스럽게 들릴 수 있지만, 꼭 그럴 필요는 없다. AI를 처음 접한다면 작게 시작하는 것을 고려해야 한다.

뉴욕시는 베테랑 소방관들에게 구조물 화재와 어떤 요인이 상관관계가 있는지에 대해 인터뷰하는 방식을 통해 AI 기반 검사 도구의 원본을 구축했다.²⁴ 불과 몇 달 만에 이러한 첫 번째 도구는 두 번째 도구를 구축하는 데 도움이 될 만큼 충분한 데이터를 수집했다. 두 도구는 이전에 스테이션 별로 분리된 카드 카탈로그와 같은 구식 시스템을 개선한 것이다.²⁵

새로운 소프트웨어를 테스트하라. 프로세스를 자동화하라. 생성형AI가 몇 가지 깔끔한 규정 준수 보고서를 안정적으로 생성할 수 있는지 확인하라. 명확하고 상당한 결과를 낳는 현실적인 문제점을 찾고 그것에 대해 평가하라. 규제 당국자라면 이미 알다시피, 때로는 시스템을 이해하려면 자신이 직접 조사를 해야 한다.



주석

1. 보잉 787 드림라이너에는 소프트웨어 코드 650만 줄이 포함되는데, 이는 물리적 부품 개수 230만 개를 훌쩍 뛰어넘는 것이다. Jeff Desjardins, "How many millions of lines of code does it take?," Visual Capitalist, February 8, 2017; Boeing, "World class supplier quality," accessed October 31, 2023.
2. Devin Coldewey, "How 'The Mandalorian' and ILM invisibly reinvented film and TV production," TechCrunch, February 21, 2020.
3. Dan Maloney, "Putting the brakes on high-frequency trading with physics," Hackaday, February 26, 2019.
4. Deloitte, "Regulatory Intelligence," homepage, accessed October 31, 2023.
5. Brian Perron, "Generative AI for social work students: Part I," Medium, March 20, 2023.
6. Angie Heise, "Generative AI and public sector," Microsoft's Public Sector Center of Expertise, accessed October 31, 2023.
7. Editorial, "AI to boost regulator in Denmark?," XBRL, February 17, 2017.
8. Alison Banney, "You can now file your tax return with help from a virtual chatbot," Finder, June 18, 2018.
9. Edward Van Buren, William D. Eggers, Tasha Austin, Joe Mariani, and Pankaj Kamleshkumar Kishnani, Scaling AI in government: How to reach the heights of enterprisewide adoption of AI, Deloitte Insights, December 13, 2021.
10. Scott W. Bauguess, "The role of big data, machine learning, and AI in assessing risks: A regulatory perspective," champagne keynote address, US Securities and Exchange Commission, June 21, 2017.
11. Issie Lapowsky, "How bots broke the Federal Communications Commission's (FCC's) public comment system," Wired, November 28, 2017; David Shepardson, "US broadband industry accused in 'fake' net neutrality comments," Reuters, May 6, 2021.
12. Jeff Kao, "More than a million pro-repeal net neutrality comments were likely faked," HackerNoon, November 23, 2017.
13. Ibid.
14. Emprata, FCC restoring internet freedom docket 17-108: Comments analysis, August 30, 2017.
15. Lorenzo Franceschi-Bicchierai, "More than 80% of all net neutrality comments were sent by bots, researchers say," Vice, October 4, 2017.
16. Deloitte, "Trustworthy AI: Bridging the ethics gap surrounding AI," accessed October 31, 2023.
17. Adam Sadilek, Stephanie Caty, Lauren DiPrete, Raed Mansour, Tom Schenk Jr., Mark Bergtholdt, Ashish Jha, Prem Ramaswami, and Evgeniy Gabrilovich, "Machine-learned epidemiology: Real-time detection of foodborne illness at scale," npj Digital Medicine 1, 2018.

18. Kristen M. Altenburger, "Artificial intelligence and food safety: Hype vs. reality," Food Safety Magazine, December 16, 2019.
19. William D. Eggers and Mike Turley, "Future of regulation: Case studies," Deloitte Center for Government Insights, accessed October 31, 2023.
20. National Science Foundation, "Fighting food poisoning in Las Vegas with machine learning," news release, March 7, 2016.
21. Jesse Roman, "In pursuit of smart," National Fire Protection Association Journal, November/December 2014.
22. William D. Eggers, Mike Turley, and Pankaj Kamleshkumar Kishnani, The regulator's new toolkit: Technologies and tactics for tomorrow's regulator, Deloitte Insights, October 18, 2018.
23. Deloitte's Center for Government Insights, "The digital citizen survey," homepage, Deloitte Insights, accessed October 31, 2023.
24. Roman, "In pursuit of smart."

저자

매튜 그레이시(Matthew Gracie)/미국

윌리엄 D. 에거스(William D. Eggers)/미국

샘 월시(Sam Walsh)/영국

스콧 스트라이너(Scott Streiner)/캐나다

판카즈 키슈나니(Pankaj Kishnani)/인도

딜로이트 산업 전문가

정부 및 공공부문(Government & Public Services), 인공지능(AI), Data Analytics 서비스

딜로이트는 정부와 공기업이 국민 삶의 질을 개선하고 국가 경쟁력을 강화할 수 있도록 지원하고 있습니다. 정부 및 공기업 부문 전문가들은 정책 제도 수립, 산업 활성화 파급효과 분석, 공공서비스의 디지털 전환 등 광범위한 분야에서 효율적이면서도 효과적으로 국가와 지역 사회를 지원하는 최적의 솔루션을 제공합니다.

정부 및 공공부문 전문팀



이재호 파트너

정부 및 공공부문 전문팀 리더 |
경영자문 부문

☎ 02 6676 2919

@ jaeholee1@deloitte.com



김정열 파트너

정부 및 공공부문 |
경영자문 부문(Deal)

☎ 02 6099 4490

@ jeongykim@deloitte.com



하성호 파트너

정부 및 공공부문 |
회계감사 부문

☎ 02 6676 1351

@ sunghha@deloitte.com



송호창 파트너

정부 및 공공부문 |
세무자문 부문

☎ 02 6676 2004

@ hochsong@deloitte.com

인공지능(AI)



조명수 파트너

디지털 경영관리 서비스 리더 |
컨설팅 부문

☎ 02 6676 2954

@ mjo@deloitte.com

Data Analytics



조민연 파트너

IT/Data Analytics |
회계감사 부문

☎ 02 6676 1990

@ minycho@deloitte.com



이성호 상무

Core Technology, Data 분석 |
컨설팅 부문

☎ 02 6676 3767

@ sholee@deloitte.com



앱



카카오톡 채널



'딜로이트 인사이트' 앱과 카카오톡 채널에서
경영·산업 트렌드를 만나보세요!



Deloitte.

Insights

성장전략부문 대표

손재호 Partner

jaehoson@deloitte.com

딜로이트 인사이트 리더

정동섭 Partner

dongjeong@deloitte.com

연구원

김사현 Director

sahekim@deloitte.com

디자이너

박근령 Senior Consultant

keunpark@deloitte.com

Contact us

krinsightsend@deloitte.com

Deloitte refers to one or more of Deloitte Touche Tohmatsu Limited (“DTTL”), its global network of member firms, and their related entities (collectively, the “Deloitte organization”). DTTL (also referred to as “Deloitte Global”) and each of its member firms and related entities are legally separate and independent entities, which cannot obligate or bind each other in respect of third parties. DTTL and each DTTL member firm and related entity is liable only for its own acts and omissions, and not those of each other. DTTL does not provide services to clients. Please see www.deloitte.com/about to learn more.

Deloitte Asia Pacific Limited is a company limited by guarantee and a member firm of DTTL. Members of Deloitte Asia Pacific Limited and their related entities, each of which are separate and independent legal entities, provide services from more than 100 cities across the region, including Auckland, Bangkok, Beijing, Hanoi, Hong Kong, Jakarta, Kuala Lumpur, Manila, Melbourne, Osaka, Seoul, Shanghai, Singapore, Sydney, Taipei and Tokyo.

This communication contains general information only, and none of Deloitte Touche Tohmatsu Limited (“DTTL”), its global network of member firms or their related entities (collectively, the “Deloitte organization”) is, by means of this communication, rendering professional advice or services. Before making any decision or taking any action that may affect your finances or your business, you should consult a qualified professional adviser.

No representations, warranties or undertakings (express or implied) are given as to the accuracy or completeness of the information in this communication, and none of DTTL, its member firms, related entities, employees or agents shall be liable or responsible for any loss or damage whatsoever arising directly or indirectly in connection with any person relying on this communication. DTTL and each of its member firms, and their related entities, are legally separate and independent entities.

본 보고서는 저작권법에 따라 보호받는 저작물로서 저작권은 딜로이트 안진회계법인(“저작권자”)에 있습니다. 본 보고서의 내용은 비영리 목적으로만 이용이 가능하고, 내용의 전부 또는 일부에 대한 상업적 활용 기타 영리목적 이용시 저작권자의 사전 허락이 필요합니다. 또한 본 보고서의 이용시, 출처를 저작권자로 명시해야 하고 저작권자의 사전 허락없이 그 내용을 변경할 수 없습니다.