



소개

루페시 초크시(Rupesh Chokshi)

애플리케이션 보안 담당 수석 부사장 겸 총괄 매니저

고객 미팅과 업계 행사, 그리고 거의 매일 뉴스를 읽을 때마다 한 가지 분명한 사실을 깨닫습니다. 새로운 AI 시대의 약속을 실현해 나가는 과정에서 직면하는 보안 문제를 인식해야 한다는 것입니다.

AI가 적절히 보호되지 않았을 때 발생하는 문제의 대표적인 사례들이 이미 나타났습니다. 악의적인 AI 조작 인시던트 중 가장 유명한 사례는 캘리포니아주 왓슨빌에서 발생했는데, 한 남성이 Chevrolet 대리점의 챗봇이 [Chevy Tahoe 신차를 1달러에 판매하도록 동의](#)하게 만든 사건이었습니다. 몇 달 후인 2024년 2월에는 [캐나다 법원이 Air Canada가 AI 기반 챗봇이 소비자에게 잘못된 정보를 제공한 것에 대해 책임이 있다고 판결](#)한 사례가 있었습니다.

물론 이러한 사례들은 초기 단계의 몇몇 사례일 뿐입니다. 현재 전 세계 기업은 자신도 모르게 새로운 AI 취약점을 자사 환경에 도입하고 있습니다. 기업 평판, 수익, 컴플라이언스 위반에 따른 벌금, 그리고 초기 AI 도입에 투자한 막대한 비용에 이르기까지 그 비용은 상당할 수 있습니다.

최근 건강검진에서 담당 의사가 AI 에이전트를 사용해 기록을 작성해도 되겠냐고 물었습니다. 그 대화는 건강 문제를 넘어 주말 계획과 제 딸의 대학 선택 등 더 넓은 주제로 이어졌습니다. 저는 그 정보가 어디로 가는지 궁금했습니다. 의사는 알고 있었을까요? 혹시 HIPAA 위반 가능성이 있었던 건 아닐까요?

이런 질문들은 전 세계 회의실과 이사회에서 제기되고 있습니다. 우리는 AI를 안전하게 사용하고 있을까요? AI를 안전하게 구축하고 있을까요? 만약 이런 질문을 하고 있지 않다면 반드시 질문해야 합니다. AI는 낙관과 혁신의 물결을 일으켰습니다. 그러나 동시에 기존 보안 솔루션으로는 감당하기 어려운 전혀 새로운 차원의 사이버 보안 취약점을 가져왔습니다. 이미 다음 두 집단 사이에서 자연스럽게 긴장이 나타나고 있습니다.

- 새로운 AI 애플리케이션과 비즈니스 모델을 서둘러 배포하려는 최고 AI 책임자와 개발팀
- 알지 못하는 위협을 어떻게 막아야 할지 고민하는 CISO