



삼일회계법인

2026년의 AI: AI 네이티브 기업

대규모 AI 도입을 위한 기업 설계





AI 네이티브 기업

비즈니스 관점에서 인공지능(AI)은 새로운 단계로 진입하고 있습니다. 파일럿 프로젝트와 알고리즘 중심의 초기 실험 단계를 지나, 이제 AI는 전략, 운영, 거버넌스, 성장 전반에 영향을 미치는 핵심 기반으로 자리 잡고 있습니다.

AI 네이티브 기업이 된다는 것은 무엇을 의미할까요?

AI는 더 이상 기존 시스템에 덧붙이는 기술이 아니라, 조직 구조와 업무 프로세스, 의사결정 방식 전반에 내재화되는 핵심 기반이 되고 있습니다. 새로운 성장 기회를 만드는 동력인 동시에, 미래 비즈니스 운영 방식을 바꾸는 요소입니다. 다만 AI 활용의 중심에는 여전히 인간의 판단과 책임, 인간 중심의 가치가 놓여야 합니다.

2026년 리더들에게 중요한 질문은 'AI를 도입할 것인가'가 아니라, AI를 중심으로 조직과 업무를 '어떻게 재설계할 것인가'입니다. 신뢰성과 속도를 확보하면서 AI의 가치를 전사적으로 확산하고 실질적인 성과로 연결하는 전략이 필요합니다.

이 보고서는 글로벌 트렌드와 주요 변화 요인을 바탕으로, 새로운 AI 환경을 형성하는 26가지 핵심 아이디어를 제시합니다. 각 아이디어는 AI 네이티브 기업이 갖추어야 할 주요 요소를 설명하며, 이를 통해 AI 네이티브 기업이 지향해야 할 방향을 보여줍니다.

목차

다음 단계	A 새로운 경쟁의 법칙	AI 기회의 전환과 새로운 경제 질서	04
전략적 우위	01 가치 창출	AI를 통한 수익 모델 재설계	07
AI를 가치 창출, 산업 경쟁력, 혁신 체계, 실행 로드맵, 성과 관리로 연결하는 방향 찾기	02 산업적 우위	산업 맥락이 만드는 경쟁력	09
	03 생태계 및 혁신	오픈 모델과 파트너십 활용	11
	04 AI 전환 로드맵	파일럿을 전사 성과로 확장	13
	05 가치 측정	성과 · 비용 · 리스크 측정	15
업무 방식의 재설계	06 리더의 역량	AI 전환을 이끄는 리더십	18
사람과 AI가 함께 일하는 방식의 전환	07 사람 중심 업무 설계	업무 흐름과 역할의 재정의	20
	08 업무 혁신	사람-AI 협업 구조로 전환	22
	09 문화 및 역량	AI 활용 문화와 역량 구축	24
지능형 시스템 구축	10 파운데이션 모델	범용 두뇌	27
AI를 현실로 만드는 기술 아키텍처	11 에이전트형 AI 시스템	자율적 행위자	29
	12 맥락(Context)	지식 및 통찰력 시스템 구축	31
	13 시뮬레이션 환경	디지털 트윈, 시뮬레이션, 연구소	33
	14 AI 운영(AI Ops)	배포 라이프라인, 모니터링, 툴링	35
설계에 의한 신뢰	15 컴퓨팅 자원 관리	전략적 자원으로서의 컴퓨팅 관리	37
AI 신뢰를 위한 설계·운영·검증 체계	16 책임 있는 AI	신뢰를 내재화하는 운영 체제	40
	17 설명 가능성	설명 가능한 의사결정 구조	42
	18 편향성, 공정성 및 유해성	공정성과 유해성의 측정 · 관리	44
	19 데이터 스튜어드십	인간 존엄성을 고려한 데이터 관리	46
	20 AI 신뢰 및 보증	상시 검증되는 신뢰 체계	48
	21 인간 중심의 통제	사람의 실질적 관리 · 감독	50
전략적 미래 사고	22 보안 및 적대적 AI	데이터 유출, 합성 위협, 에이전트 리스크	53
리더가 지금 예측해야 할 새로운 위험	23 안전 및 시스템 리스크	시스템 불안정성과 모델 붕괴 방지	55
	24 국가 인프라	AI 인프라 기반과 제약 요인	57
	25 제로 이미지션 인텔리전스	AI 환경 발자국 관리	59
	26 미래예측 거버넌스	장기적 AI 변화에 대비하는 거버넌스	61
다음 단계의 리더십	Ω AI 확장을 위한 핵심 역량	끊임없이 변화하는 AI 미래를 자신 있게 이끄는 방법	63
AI 활용 단계	- AI의 수준 – 생성형 AI, 에이전트 및 기타 2025년 최신 기술 및 2026년 신호 - AI 용어집		67 68 70
본 보고서에 대하여			71

다음 단계

AI의 잠재력, 수조 달러 규모 산업의 부상, 자율 에이전트와 로봇의 현실화, 대형 딥, 데이터센터 인프라 투자, 규제 강화, AI 전문 인력 확보를 위한 파격적인 처우 등 연일 이어지는 헤드라인은 AI가 산업 전반에 지각변동을 일으키고 있음을 보여줍니다. AI는 이미 기술 기반의 차세대 GDP 성장을 견인할 잠재력을 입증하고 있습니다.

이제 리더들이 던져야 할 질문은 AI의 가능성이 아니라, 이를 성과로 전환하기 위한 실행입니다. “2026년, 지금 무엇을 실행해야 하는가?”

실험에서 실전으로

AI는 이제 실험 단계를 넘어 기업 경영의 핵심 의제로 부상했습니다. 정책적 지원, 기업 투자, 산업적 연계, 혁신 역량이 맞물리면서 AI 확산을 위한 기반도 빠르게 갖춰지고 있습니다. 수년간 파일럿과 데모를 거친 AI는 더 이상 호기심의 대상이나 부가적 프로젝트가 아니라, 기업의 시스템과 프로세스, 의사결정을 관통하는 필수 인프라로 자리 잡고 있습니다.

이에 따라 리더들의 관심은 AI 도입 여부가 아니라 전사적 확산과 성과 창출에 맞춰지고 있습니다. 이제 중요한 것은 AI를 실제 업무와 프로세스에 안착시키고, 이를 조직 차원의 성과로 확장하는 것입니다.

CEO의 인식 변화

오늘날 CEO들은 AI를 단순한 기술 도입 과제가 아니라 전략과 운영을 뒷받침하는 핵심 기반으로 인식하고 있습니다. PwC의 2025년 CEO 설문조사에 따르면, 많은 CEO들이 생성형 AI를 효율성 개선과 수익성 제고, 매출 성장에 영향을 미치는 주요 요인으로 보고 있습니다.

AI 도입 효과는 이미 일부 기업에서 나타나고 있지만, 성과는 모든 기업에 고르게 나타나지 않습니다. 단편적인 시범 적용에 머문 기업은 비용 부담만 키우는 반면, 명확한 목표와 실행 체계를 갖춘 기업은 AI를 생산성 향상과 새로운 사업 기회로 연결하고 있습니다.

따라서 AI 확산은 개별 프로젝트 차원이 아니라 전사적 실행 과제로 다뤄져야 합니다. AI를 전략과 연결하고, 실제 업무와 프로세스에 안착시키며, 신뢰와 성과 기준에 따라 관리하는 것. 이것이 ‘AI 네이티브’ 기업으로 나아가기 위한 출발점입니다.

새로운 경쟁의 법칙

AI는 미래 비즈니스의 핵심 기반으로 자리 잡으며, 기업의 기회 포착 방식과 경쟁 구도를 바꾸고 있습니다. 고객 서비스에서는 단순한 챗봇 도입을 넘어, 고객 문의를 이해하고 답변 초안을 작성하며 필요한 경우 상담사에게 연결하는 방식으로 AI가 업무 전반에 적용되고 있습니다. 제조업에서도 일부 설비의 예측 유지보수에 그치지 않고, 설비 상태를 실시간으로 파악하고 이상 징후에 대응하는 체계로 전환되고 있습니다.

이러한 변화는 AI를 개별 업무의 효율화 수단이 아니라, 기업 경쟁력의 핵심 조건으로 재정의하고 있습니다. 앞으로 기업의 경쟁 우위는 시장 대응 속도, 확장 가능한 운영 체계, 지속적인 혁신 역량, 신뢰 기반의 의사결정 능력에 따라 좌우될 것입니다.

- **속도:** 선도 기업은 제한적인 실험에 머물지 않고, 핵심 업무에 AI를 빠르게 적용해 시장 변화에 대응하고 있습니다.
- **규모:** AI는 인력과 조직 규모에 의존하던 기존 성장 방식을 바꾸고, 제한된 자원으로도 운영 범위를 확장할 수 있게 합니다.
- **혁신:** AI는 새로운 제품과 서비스, 비즈니스 모델을 빠르게 검증하고 확산하는 기반이 됩니다.
- **신뢰:** 책임 있는 AI 체계는 도입 과정에서 발생할 수 있는 리스크를 관리하고, 의사결정의 안정성을 높입니다.

2026년의 AI

이 보고서는 2026년 AI 환경을 형성할 26가지 핵심 아이디어를 정리한 것입니다. 각 아이디어는 AI 기반 기업을 구성하는 핵심 영역을 보여주며, 서로 독립적인 주제이면서도 유기적으로 연결되어 있습니다.

PwC의 글로벌 사례와 연구 결과를 바탕으로, 기업이 AI 네이티브 조직으로 전환하기 위해 고려해야 할 전략적 방향을 제시합니다.

2026년 진입 시점

PwC 제28·29차 글로벌 CEO 설문조사 기준

- AI를 핵심 경영 전략으로 인식: 글로벌 CEO **90%**
- AI 투자 수준 충분: 글로벌 CEO 40%
- AI 도입을 뒷받침하는 조직문화 보유: **66%**
- AI 운영 기반 확보: **18%**

실험 단계

AI 활용 가능성을 검증하는 초기 단계입니다. 개념증명 (PoC), 파일럿, 개별 혁신 과제를 통해 리스크가 낮은 업무나 특정 부서에서 적용 가능성을 확인합니다.

다만 대부분의 과제가 핵심 업무와 전사 시스템에서 분리되어 있어 확장에는 한계가 있습니다. 데이터 품질, 기술 구조, 거버넌스, 변화관리 체계도 아직 정비되지 않아 AI는 조직 운영을 바꾸는 기반보다 개별 업무 도구로 활용되는 수준에 머무릅니다.

결과: 부분적 성과는 나타나지만 전사적 변화로 이어지지 않음

AI 네이티브 기업 운영 체계

AI가 실험적 도구를 넘어 기업 운영과 의사결정 구조에 내재화되는 단계입니다. AI는 기존 업무에 추가되는 기능이 아니라, 처음부터 업무 체계 안에 설계됩니다.

리더십의 AI 이해도, 전사적 조율, 기술 통제 기준이 핵심 조건이 됩니다. AI는 신뢰를 전제로 설계되고, 안정적으로 확장되며, 속도·정확성·경제성·책임성을 높이는 운영 체계로 자리 잡습니다.

결과: AI가 기업 운영의 기본 방식으로 정착



구분	실험 단계	AI 네이티브 기업 운영 체계
전략	개별 과제 중심	기업 가치 창출 중심
경쟁	경쟁사 수준 대응	차별화 수단으로 활용
리더십	기술 조직 중심	리더십 전반의 AI 전환 주도
가치 실현	파일럿 성과 중심	사업 가치 창출 영역 관리
업무 프로세스	기존 업무에 일부 적용	AI를 전제로 업무 재설계
역량	개인 학습 중심	사람-AI 협업 역량 확산
문화	일부 주도 인력 중심	AI 활용 문화 정착
인력	기술 인재 채용 중심	사람-AI 역할 기반 인력 계획
운영 모델	사안별 대응	부서 간 협업 기반 운영
기술 구조	개별 도구 활용	데이터·플랫폼 통합 관리
신뢰와 리스크	규제 대응 및 사후 통제	설계 단계부터 신뢰·리스크 관리 반영

전략적 우위

전략적 우위는 AI를 수익 모델, 산업 경쟁력, 혁신 방식, 운영 체계, 성과 관리 기준으로 연결하는 과정입니다.

AI는 기업의 전략과 시장 대응 방식을 바꾸고 있습니다. 가치 창출은 AI를 수익 모델과 투자 판단에 연결하고, 산업적 우위는 산업 전문성과 고유 데이터를 경쟁력으로 전환합니다. 생태계와 혁신은 외부 기술과 파트너십을 활용해 속도와 선택권을 높이며, AI 전환 로드맵은 일회성 파일럿을 전사적 실행 체계로 확장합니다. 가치 측정은 수익, 비용, 리스크, 고객 경험에 미치는 영향을 데이터로 확인해 AI 확장과 투자 판단을 뒷받침합니다.



진정한 가치 창출 전략의 핵심은 AI가 비즈니스 모델의 손익(P&L) 구조와 시장 전반의 가치 창출 영역(Value Pool)을 어떻게 재편하는지 파악하는 데 있습니다.

가치 창출은 기존 프로세스에 AI를 단순히 추가하는 일이 아닙니다. 제품 기획, 가격 책정과 제공 방식, 고객 확보와 응대, 리스크 관리, 사업 영역 설정에 이르기까지 기업이 수익을 만들어내는 방식을 다시 설계하는 과정입니다.

국내 기업의 성장도 더 이상 제조 역량과 수출 규모 확대만으로 설명되기 어렵습니다. 앞으로의 경쟁력은 고객과 시장의 수요를 얼마나 정교하게 파악하고, 이를 새로운 제품과 서비스, 수익 모델로 연결하는 능력에 달려 있습니다. AI는 원가 구조를 개선하고 제품·서비스의 범위를 확장하며, 새로운 수익 기회를 발굴하는 데 활용될 수 있습니다. 이에 따라 가치가 창출되는 영역도 빠르게 달라지고 있습니다.

관건은 이익 증가가 어디에서 발생하는지 명확히 파악하고, AI를 업무 효율화와 제품 경쟁력 강화에 실질적으로 연결하는 것입니다.

가치 창출은 전사적 관점에서 접근하되, 단기 성과와 장기 성장 과제를 함께 관리해야 합니다. 단기적으로는 고객 전환율 개선, 이탈률 감소, 처리 속도 단축, 리스크 판단 고도화 등 고객 접점과 운영 과정에서 수익성을 높일 수 있는 영역에 집중해야 합니다. 동시에 데이터 기반 서비스, 구독형 서비스, 전략적 파트너십 등 새로운 성장 영역에 대한 투자도 병행해야 합니다.

핵심은 AI 성과가 어디에서 발생했는지를 정교하게 측정하고, 이를 다음 혁신과 투자로 연결하는 체계를 갖추는 것입니다. 성과를 빠르게 입증하고 확산의 근거로 삼을 때, AI의 잠재력은 실제 수익으로 이어질 수 있습니다.

왜 지금 중요한가

2026년까지 AI로 명확한 투자 수익률(ROI)을 확보하는 일이 최우선 과제입니다. 수년간의 실험을 거치며 2026년까지 기업의 핵심 과제는 AI 투자에서 명확한 투자수익률(ROI)을 확보하는 것입니다. 수년간의 실험을 거치며 기업은 AI 투자가 기술 시연을 넘어 지속 가능한 수익성으로 이어지기를 기대하고 있습니다. 그러나 여전히 상당수 AI 이니셔티브는 실질적 성과를 충분히 입증하지 못하고 있습니다. 따라서 혁신의 초점은 '개념 증명(PoC)'에서 '가치 증명(PoV)'으로 이동해야 하며, 확장과 내재화를 뒷받침할 새로운 사업 운영 원칙이 필요합니다.

2026년에 변화하는 것

기업의 가치 창출 방식은 다음과 같이 변화하고 있습니다:

- **지표의 확장:** 비용 절감과 효율성 중심의 지표를 넘어, AI가 매출 성장, 고객 경험, 혁신에 미치는 영향을 함께 측정합니다.
- **성과 창출 다각화:** 초기 AI 활용이 업무 자동화와 생산성 향상에 집중됐다면, 이제는 제품 기획과 기능 고도화를 통해 매출을 늘리고 고객생애가치(LTV)를 높이는 방향으로 확장됩니다.
- **새로운 가치 흐름:** 선도 기업은 데이터 기반 구독 서비스와 스마트 제품 등 AI 기반의 새로운 서비스와 비즈니스 모델을 통해 수익원을 다변화하고 있습니다. 이에 따라 가치 창출의 범위는 단순한 비용 절감을 넘어 비즈니스 모델 재설계로 확대됩니다.
- **AI 기반 기회 발굴:** AI의 심층 분석 역량을 활용해 시장과 기업 내부에서 잠재 기회의 신호를 지속적으로 포착하고, 이를 사업 기회로 연결합니다.

혁신 방향은 ‘개념 증명(PoC)’에서 ‘가치 증명(Proof of Value)’으로 전환

- 개별 성과에 머무르지 않고, 기업 전체의 가치 창출 영역을 최적화하는 하향식 관점으로 전환
- 단기 운영 효율화와 장기 사업 성장 전략을 병행하는 이중 트랙(dual-track) 가치 창출 체계 정립
- AI 투자의 실질적 성과를 확인하기 위한 수익과 투자자본수익률(ROIC)의 가시성 확보

한국의 상황

국내 기업은 높은 디지털 도입 수준과 축적된 데이터, 다양한 고객 접점을 바탕으로 AI를 수익 창출 역량으로 전환할 수 있는 기반을 갖추고 있습니다. 특히 금융, 유통, 헬스케어, 에너지, 반도체·제조 분야에서는 가격 최적화, 서비스 개인화, 운영 효율화가 AI의 주요 활용 영역으로 자리 잡고 있습니다.

앞으로의 관건은 AI를 개별 기능이나 실험 과제에 머물게 하지 않고, 사업 운영 전반에 내재화하는 것입니다. 국내 데이터와 시장 이해를 바탕으로 가격 책정을 정교화하고, 고객 서비스를 개인화하며, 인접 사업 영역으로 확장하는 역량은 기업의 경쟁 우위를 좌우하는 핵심 요소가 될 것입니다.

리더십 우선순위

- **AI와 수익의 연계:** AI가 매출이나 이익에 직접 영향을 줄 수 있는 고객 여정과 제품을 선정하고, 책임자와 평가 기준을 명확히 합니다.
- **성과 흐름의 계측:** AI의 가치 기여도를 확인하고 논의할 수 있도록 업무 프로세스 안에 기여도 측정 체계를 구축합니다. 대조군 설정, 사업 KPI 도출, 실험 설계와 검증 절차가 이에 포함됩니다.
- **자본의 집중 배분:** 다수의 소규모 PoC에서 벗어나, 명확한 투자 회수 계획을 갖춘 확산형 포트폴리오로 자금을 이동합니다.
- **CFO와의 조율:** AI 도입이 매출, 비용, 운영 구조에 미치는 영향을 손익(P&L)에 반영하고, 지속 운영 비용과 회복탄력성, 성과 추적 계획을 함께 수립합니다.

변화의 신호

- **가치 기여도 측정:** 월간 검토를 통해 AI 활용 과제와 사업 KPI의 연계성을 점검하고, 성과가 낮은 과제는 개선하거나 중단합니다.
- **운영 내재화:** AI를 독립된 실험 과제가 아니라 고객 접점, 제품 개발, 운영 프로세스에 통합합니다.
- **경제성 개선:** 서비스 제공 비용을 낮추고 고객 만족도와 이용 지표를 개선함으로써, AI 기반 서비스의 수익성을 높입니다.
- **성과 기반 재투자:** 비용 절감 효과와 신규 수익을 제품 고도화와 업무 프로세스 개선에 재투입합니다.

AI 리더의 사고방식

AI를 비용 항목이 아니라 가치 창출 엔진으로 대합니다.

AI가 적용되는 고객 접점, 제품·서비스, 업무 프로세스별로 손익 기준을 설정하고, 성과와 경제성에 따라 자본을 배분합니다. 창출된 수익은 지속적인 제품 고도화와 운영 방식 재설계에 재투입하며, 이를 통해 AI가 비즈니스 모델 전반의 변화를 이끌도록 합니다.

즉, AI 기능과 서비스를 내장해 새로운 수익 곡선과 이익을 만들어내는 동시에, 측정과 책임을 운영의 핵심에 둡니다. 선도적 리더는 AI를 전략의 필수 요소로 보고, “AI가 핵심 수익과 고객 지표를 어떻게 움직이고 있는가, 다음 전략 단계는 무엇인가”를 지속적으로 점검합니다.

범용 AI만으로는 경쟁 우위가 만들어지지 않습니다. 산업 이해, 고유 데이터, 업무 방식, 규제 대응 역량이 결합될 때 차별화된 성과가 만들어집니다.

2026년 AI 경쟁력은 기업이 속한 산업 현장에서 결정됩니다. 더 큰 모델을 보유하는 것보다 중요한 것은 산업별 업무 흐름과 리스크 관리 기준, 규제 요건, 고객 특성을 AI에 얼마나 깊이 반영하느냐입니다.

경쟁사가 유사한 AI 기능을 쉽게 구현할 수 있다면, 그것은 아직 차별적 우위가 아닙니다. 자사 데이터, 업무 프로세스, 의사결정 기준, 검증 체계가 결합되어야 모방하기 어려운 경쟁력이 만들어집니다. AI가 산업별 현장 지식과 운영 방식에 깊이 연결될수록 경쟁 우위는 더 분명해집니다.

산업적 우위는 단순히 AI 기술을 도입하는 데서 나오지 않습니다. 선도 기업은 자사 산업에서 AI가 실질적 성과를 낼 수 있는 영역을 선별하고, 검증된 사례를 자사 사업 맥락에 맞게 적용합니다. 이를 통해 고객 대응, 가격 책정, 리스크 판단, 운영 효율화 등 핵심 업무에서 AI 활용 수준을 높여갑니다.

중요한 것은 속도에 대한 기준도 달라지고 있다는 점입니다. 같은 기반 모델을 활용하더라도, 신뢰할 수 있는 AI를 더 빠르게 업무에 적용하는 기업이 앞서갑니다. 규제 준수, 출처 확인, 설명 가능성, 안전장치가 설계 단계부터 반영될 때 AI의 도입 속도는 곧 경쟁력이 됩니다.

결국 산업적 우위는 AI를 자사 산업의 문제 해결 방식에 맞게 조정하는 데서 나옵니다. 표준화된 도구를 그대로 도입하는 것만으로는 충분하지 않습니다. 산업 규제, 고객 기대, 현장 데이터, 전문 지식이 결합될 때 경쟁사가 쉽게 따라오기 어려운 차별성이 만들어집니다.

왜 지금 중요한가

경영진의 관심은 이제 AI 도입 여부를 넘어, AI가 산업 내 경쟁 구도와 자사의 위치에 미치는 영향으로 옮겨가고 있습니다. AI 투자는 기술 실험이 아니라 핵심 사업 지표 개선으로 연결되어야 합니다.

은행의 수익성, 에너지 사업자의 운영 안정성, 의료 분야의 환자 결과, 제조업의 품질과 생산성처럼 산업마다 중요한 성과 지표는 다릅니다. 따라서 AI도 각 산업의 언어로 성과를 입증해야 합니다. 앞으로의 과제는 AI를 실험 단계에서 산업별 성과를 만들어내는 실행 역량으로 전환하는 것입니다.

2026년에 변화하는 것

AI 도입은 산업별 요구와 규제 환경에 더 깊이 맞춰지는 방향으로 전개되고 있습니다.

- **규제의 본격화:** 2026년 1월 「AI 기본법」 시행을 계기로 AI의 신뢰성, 안전성, 투명성 확보가 기업의 핵심 과제로 부상했습니다. 신뢰할 수 있는 AI를 빠르게 구현하는 역량이 시장 경쟁력으로 연결됩니다.
- **산업별 AI 솔루션 확산:** 제조, 금융, 의료 등 주요 산업에 특화된 AI 플랫폼과 서비스가 확대되고 있습니다. 산업별 데이터와 업무 흐름을 반영한 AI가 실제 성과 창출의 중심이 됩니다.
- **산업 전문성과 AI 역량의 결합:** AI 기술만 이해하는 인력보다, 산업 전문성과 AI 활용 역량을 함께 갖춘 인재와 조직의 중요성이 커집니다. 사람, 프로세스, 기술이 만나는 지점에서 실질적인 사업 가치가 만들어집니다.
- **소버린 AI와 독자 모델 강화:** 오픈 모델을 국내 데이터로 조정하고, 국내 규제와 산업 특성에 맞는 안전장치를 반영하는 흐름이 확대됩니다. 정부의 독자 AI 파운데이션 모델 사업과 맞물려 데이터 주권과 기술 자립의 중요성도 커지고 있습니다.

범용 AI는 경쟁 우위를 만들지 못하지만, 산업의 맥락은 경쟁 우위를 만듭니다.

글로벌 AI 기술은 누구나 활용할 수 있습니다.

그러나 그 가치는 데이터, 업무 흐름, 전문성, 신뢰 체계를 결합해
자사 산업의 문제를 해결하는 기업에게 돌아갑니다.

한국의 상황

국내 기업은 엄격한 규제 체계와 산업별 특성이 뚜렷한 환경에서 사업을 운영하고 있습니다.

금융 · 의료처럼 규제 강도가 높은 산업과 반도체 · 제조처럼 공정 노하우가 핵심인 산업이 공존합니다. 특히 반도체를 중심으로 한 수출 구조는 산업별 맞춤형 AI 역량의 중요성을 더욱 높이고 있습니다.

또한 데이터 주권과 국내산업 특성을 반영한 소버린 AI 전략이 정부 사업을 중심으로 본격화되고 있습니다. 이는 국내 기업이 글로벌 AI 기술을 활용하는 동시에, 자사 데이터와 산업 지식을 기반으로 독자적 경쟁력을 확보해야 한다는 점을 보여줍니다.

리더십 우선순위

- **핵심 전문 역량에 우선 투자:** 글로벌 AI 모델을 활용하되, 자사 산업 데이터와 업무 노하우를 결합해 차별화된 경쟁력을 구축해야 합니다.
- **규제 대응의 선제적 내재화:** 설계 단계부터 '고위험 인공지능(AI)' 해당 여부를 점검하고, 위험관리 체계와 표시 의무 등 법적 요건을 업무 프로세스에 반영해야 합니다.
- **고유 데이터의 전략적 자산화:** 자사의 고유 데이터를 식별하고 정비해 AI 경쟁력의 원천으로 전환해야 합니다. 쉽게 대체할 수 없는 데이터 자산이 산업적 우위의 기반이 됩니다.

변화의 신호

- **핵심 업무에 AI 내재화:** AI가 파일럿 단계를 넘어 신용평가, 수요 예측, 생산 공정 관리 등 핵심 업무에 적용됩니다.
- **규제 신뢰성 확보:** 외부 검토와 내부 점검에서 AI 기반 업무 프로세스의 안전성, 공정성, 설명 가능성이 입증됩니다.
- **차별화된 고객 가치:** AI 기반 기능이 고객 경험과 서비스 품질을 개선하고, 시장에서 분명한 차별화 요소로 작용합니다.
- **산업 데이터의 축적과 활용:** AI 시스템이 운영 과정에서 데이터를 지속적으로 축적하고, 이를 다시 성과 개선에 활용하는 선순환 구조가 형성됩니다.

AI 리더의 사고방식

리더는 AI를 범용 기술이 아니라 자사 산업의 문제를 해결하는 실행 수단으로 바라봅니다.

중요한 것은 AI 기술 자체가 아니라, 이 기술이 우리 산업의 어떤 문제를 해결하는지, 고객과 규제의 요구를 충족하는지, 실제 사업 성과로 연결되는지를 지속적으로 점검하는 것입니다. AI가 산업 맥락과 결합될 때, 범용 기술은 차별적 경쟁력으로 전환됩니다.

AI 시대의 경쟁력은 개별 기업의 역량만으로 완성되지 않습니다.

혁신은 연결된 생태계 안에서 더 빠르게 확장됩니다.

외부 기술과 파트너십, 산업 전문성을 결합하는 역량이 성과의 차이를 만듭니다.

AI 경쟁력은 더 이상 개별 기업이 보유한 기술만으로 결정되지 않습니다. 2026년의 선도 기업은 모든 AI 역량을 내부에서 직접 구축하기보다, 오픈 모델과 전문 솔루션, 클라우드 인프라, 파트너사의 기술을 목적에 맞게 조합해 AI 전략을 추진하고 있습니다. AI 기술이 빠르게 변화하는 환경에서는 모든 역량을 자체적으로 확보하고 운영하는 방식이 비용과 속도 면에서 한계를 드러낼 수밖에 없습니다.

중요한 것은 필요한 기술과 파트너를 적시에 선택하고, 이를 자사 데이터와 업무 맥락에 맞게 결합하는 능력입니다. 기업들은 대형 클라우드 사업자의 인프라와 AI 서비스를 활용하고, 오픈 모델을 통해 유연성과 비용 효율성을 확보하며, 학계와 스타트업의 전문성을 접목해 새로운 기술과 아이디어를 빠르게 도입하고 있습니다. 이처럼 내부 역량과 외부 혁신을 연결하는 방식은 AI 도입의 속도와 확장성을 높이는 핵심 수단이 되고 있습니다.

생태계 전략은 혁신 속도와 시장 대응력을 동시에 높입니다. 새로운 모델과 서비스가 계속 등장하는 상황에서 경쟁 우위는 특정 기술을 독점하는 데서 나오기보다, 적합한 기술을 빠르게 검증하고 사업에 적용하는 역량에서 만들어집니다. 다양한 기술과 공급자를 활용하는 기업은 기술 변화와 비용 구조 변화에도 유연하게 대응할 수 있으며, 새로운 사업 기회도 더 빠르게 포착할 수 있습니다.

따라서 선도 기업은 AI를 단일 플랫폼이나 특정 공급자의 기술로만 바라보지 않습니다. AI를 기술, 서비스, 파트너십이 결합된 확장형 생태계로 인식하고, 내부 역량과 외부 혁신을 전략적으로 연결하는 데 집중합니다. AI 발전 속도가 빨라질수록 경쟁력은 개별 기술 자체보다, 어떤 생태계를 구축하고 이를 얼마나 효과적으로 활용 하느냐에 따라 결정될 것입니다.

왜 지금 중요한가

2026년 AI 협력은 임시 대응이 아니라 명확한 운영 원칙을 전제로 합니다. 외부 파트너나 오픈 모델을 활용할 때는 성과 지표, 데이터 출처, 비용 구조, 데이터와 지식재산의 공유 범위를 사전에 정해야 합니다. 이러한 기준은 기술 변화와 비용 변동, 규제 리스크에 대응하면서 특정 공급자나 플랫폼에 대한 과도한 의존을 줄이는 기반이 됩니다.

협력 체계는 글로벌 수준의 안정성과 국내 시장 이해를 함께 갖춘 파트너십을 기반으로 설계되어야 합니다. 기업은 파트너가 제공하는 성과와 리스크를 정기적으로 점검하고, 보안 · 출처 · 책임 기준을 충족하는 협력 구조를 마련해야 합니다. 이러한 구조가 갖춰질 때 외부 협업은 혁신 속도를 높이는 동시에 자사 데이터와 운영 노하우를 보호하는 수단이 됩니다.

2026년에 변화하는 것

AI 생태계 전반에 새로운 운영 방식이 요구되고 있습니다:

- **오픈 모델의 확산:** 기업용 오픈 모델이 선택지를 넓히고 특정 공급자 의존도를 낮추고 있습니다. 성능뿐 아니라 보안, 출처, 편향성 검증이 주요 기준으로 부상하고 있습니다.
- **협력 네트워크 확대:** 산업 간 컨소시엄과 파트너십이 공동 과제 해결과 실무 표준 마련의 기반이 되고 있습니다.
- **모듈형 AI 구조 정착:** 기술과 비용 환경 변화에 맞춰 모델, 인프라, 서비스 구성 요소를 유연하게 교체하는 방식이 확산되고 있습니다.
- **생태계 거버넌스 강화:** 데이터, 지식재산, 보안, 윤리 기준을 명확히 하는 것이 외부 협업의 핵심 조건이 되고 있습니다.

어떤 기업도 AI 시대를 홀로 이끌 수는 없습니다.

혁신

AI 시대의 혁신은 R&D 부서나 초기 실험에 머무르지 않습니다. 제품 개발, 서비스 개선, 고객 대응, 운영 효율화 전반에 스며드는 실행 역량이 되고 있습니다.

AI 혁신은 더 빠른 의사결정, 개인화된 경험, 정교한 자동화를 통해 새로운 수익원과 성장 동력을 찾는 과정입니다. 중요한 것은 속도만이 아니라 방향입니다. 혁신은 비즈니스 목표와 연결되어야 하며, 핵심 운영 흐름 안에서 책임과 예산, 성과 기준에 따라 관리되어야 합니다.

검증된 원칙은 여전히 유효합니다. 혁신은 기업의 가치 창출 영역에서 출발해야 하며, 실제 사업 성과로 입증되고 확장될 수 있어야 합니다. 부서 간 협업을 기반으로 지속적인 혁신 과정을 운영하고, 성과 측정과 피드백을 통해 반복적으로 개선하는 구조가 필요합니다.

한국의 상황

국내 AI 생태계는 우수한 대학·연구기관, 기술 스타트업, 대기업, 글로벌 플랫폼이 결합하는 구조로 발전하고 있습니다. 기업들은 국내 스타트업의 특화 기술과 글로벌 플랫폼의 확장성을 함께 활용하고 있으며, 정부의 독자 AI 파운데이션 모델 사업을 계기로 산학연·대기업·스타트업 컨소시엄도 확대되고 있습니다.

다만 개방형 협력에는 명확한 관리 기준이 필요합니다. 오픈 모델과 외부 파트너 활용은 선택지를 넓히지만, 개인정보보호, 클라우드 보안, 데이터 관리 기준이 함께 작동해야 합니다. 핵심 데이터와 지식재산이 외부 협업 과정에서 희석되지 않도록 통제 체계를 갖추는 것이 중요합니다.

리더십 우선순위

- **전략적 파트너십:** 자사의 AI 목표에 맞는 핵심 파트너를 선별하고, 역할과 책임, 성과 기준을 명확히 해야 합니다.
- **협업 거버넌스:** 데이터 사용, 지식재산 권리, 보안 기준, 계약 종료 조건 등 핵심 규칙을 사전에 정해 협업 리스크를 관리해야 합니다.
- **유연한 기술 구조:** 특정 기술이나 공급자에 종속되지 않도록 시스템과 계약 구조를 설계하고, 기술 선택의 유연성을 유지해야 합니다.
- **내·외부 역량 결합:** 내부 팀과 외부 파트너가 함께 개발·검증에 참여할 수 있는 체계를 마련해 속도와 신뢰성을 동시에 확보해야 합니다.

변화의 신호

- **협력 네트워크 운영:** 외부 전문가와 내부 팀의 공동 개발이 실제 서비스 출시와 운영으로 이어지고 있습니다.
- **다중 모델 활용 확대:** 다양한 출처의 모델과 서비스를 업무 목적에 맞게 조합해 활용하는 방식이 확산되고 있습니다.
- **혁신 주기 단축:** 검증된 오픈 모델과 파트너 API를 활용해 새로운 요구를 짧은 주기로 운영 시스템에 반영하는 사례가 늘고 있습니다.
- **협업 성과 가시화:** 외부 협력이 단순 실험을 넘어 수익 증대와 비용 효율 개선으로 연결되고 있습니다.

AI 리더의 사고방식

2026년의 AI 리더는 기술의 소유자가 아니라 생태계의 조율자입니다. 내부 역량과 외부 파트너십을 연결해 혁신 속도를 높이되, 데이터·지식재산·리스크에 대한 통제력은 유지해야 합니다.

이제 핵심 질문은 “우리가 무엇을 직접 만들 수 있는가”에 머물지 않습니다. “어떤 역량과 연결하고, 누구와 협력해야 더 빠르게 성과를 만들 수 있는가”가 AI 리더십의 중요한 기준이 되고 있습니다.

AI 전환 로드맵은 일회성 파일럿을 전사적 성과로 확장하는 실행 기준입니다.
핵심 업무와 운영 구조에 AI를 단계적으로 내재화해 반복 가능한 성과 창출 체계를 구축합니다.

AI 전환 로드맵은 기업이 AI와 함께 일하는 방식을 체계적으로 정비하기 위한 기준입니다. 전략, 운영, 데이터, 거버넌스, 기술 구조를 하나로 연결하고, 초기 성과가 일부 팀이나 개별 프로젝트에 머물지 않도록 전사 차원의 실행 체계로 확장하는 데 목적이 있습니다.

대부분의 조직은 특정 업무나 기능에 AI를 먼저 적용하는 방식으로 출발합니다. 그러나 AI가 핵심 업무와 의사결정 과정으로 확대되면, 단순한 실험 관리만으로는 충분하지 않습니다. 전략과 운영 모델, 데이터 기준, 기술 구조, 거버넌스가 함께 정렬되어야 AI가 조직 전체의 성과 개선으로 이어질 수 있습니다.

좋은 로드맵은 복잡한 계획서가 아닙니다. 어디에 AI를 적용할지, 어떤 기준으로 성과를 판단할지, 어떤 데이터와 기술 구조가 필요한지를 명확히 정리한 실행 기준입니다. 이를 통해 각 부서는 같은 방향으로 움직이고, 중복 투자를 줄이며, 검증된 성과를 빠르게 확산할 수 있습니다.

AI 전환 로드맵은 기존 업무 방식을 그대로 자동화하는 데 그치지 않습니다. AI를 전제로 업무 절차와 의사결정 구조를 다시 설계하고, 필요한 역할과 책임을 명확히 하는 과정입니다. 중요한 것은 기존 시스템을 무리하게 바꾸는 것이 아니라, AI가 실질적인 변화를 만들 수 있는 영역부터 단계적으로 전환하는 것입니다.

왜 지금 중요한가

많은 기업이 이미 수십 개의 AI 파일럿과 도구를 운영하고 있지만, 이를 전사적으로 확장할 통합된 방법은 많은 기업이 이미 여러 AI 파일럿과 도구를 운영하고 있지만, 이를 전사적으로 확장하는 방식은 아직 충분히 정착되지 않았습니다. 개별 프로젝트의 성과가 조직 전체의 생산성과 수익성 개선으로 연결되지 못하는 이유도 여기에 있습니다.

2026년에는 AI를 혁신 실험 단계에서 실제 운영 자산으로 전환하는 역량이 핵심 과제가 됩니다. AI가 전략 과제와 연결되고, 업무 프로세스 안에 내재화되며, 데이터와 거버넌스 체계 아래에서 관리될 때 성과는 반복적으로 만들어질 수 있습니다.

따라서 AI 도입은 단순한 기술 적용이 아니라 전사 운영 모델의 변화로 봐야 합니다. 어떤 업무를 다시 설계할지, 어떤 의사결정에 AI를 활용할지, 성과와 리스크를 어떻게 관리할지에 대한 명확한 실행 로드맵이 필요합니다.

2026년에 변화하는 것

AI 전환 로드맵은 더 이상 고정된 계획서가 아니라, 지속적으로 조정되는 운영 체계로 바뀌고 있습니다:

- **제로베이스 검토:** 기존 업무 절차를 그대로 자동화하는 것이 아니라, AI를 전제로 업무와 의사결정 방식을 다시 설계합니다.
- **통합 실행 계획:** AI 전략, 데이터, 기술 구조, 거버넌스를 하나의 실행 계획으로 연결해 부서별 도입의 단절을 줄입니다.
- **업무별 적용 기준 정립:** 어떤 업무에 AI를 적용할지, 어떤 성과 지표로 판단할지, 어떤 위험을 관리해야 하는지를 사전에 정의합니다.
- **확장 가능한 기술 구조:** 파일럿 단계의 도구가 아니라 전사 운영 환경에서 재사용 · 확장 가능한 모델, 데이터, 시스템 구조를 마련합니다.

제로베이스 검토는 AI가 점진적 개선을 넘어 어디에서 구조적 변화를 만들 수 있는지 파악하게 합니다.

이는 리더에게 다음과 같은 질문을 던집니다.

“AI를 전제로 이 업무를 처음부터 다시 설계한다면
무엇을 다르게 할 것인가?”

한국의 상황

국내 기업은 전사 차원의 AI 도입 필요성을 빠르게 인식하고 있습니다. 특히 금융, 제조, 유통, 통신 등 주요 산업에서는 AI를 개별 기능이나 부서 단위의 실험이 아니라 전사 운영 혁신의 과제로 다루기 시작했습니다.

2026년 「AI 기본법」 시행을 계기로 전사 운영 체계의 중요성도 커지고 있습니다. AI의 신뢰성, 안전성, 투명성, 위험관리 기준을 설계 단계부터 반영해야 하기 때문입니다. 이에 따라 국내 기업은 AI 도입을 기술 과제가 아니라 전략, 데이터, 거버넌스, 운영 체계를 함께 정비하는 과제로 바라봐야 합니다.

특히 최고재무책임자(CFO), 최고정보책임자(CIO), 최고데이터책임자(CDO), 법무·리스크 조직의 역할이 중요해지고 있습니다. AI 투자와 성과, 데이터 관리, 규제 대응, 운영 리스크가 하나의 전사 관리 체계 안에서 다뤄져야 하기 때문입니다.

리더십 우선순위

- **로드맵 수립:** 전략, 운영, 데이터, 기술, 거버넌스를 연결하는 전사 AI 실행 기준을 마련합니다.
- **거버넌스 구축:** 표준과 책임 구조를 정의하고, AI 리스크를 관리할 수 있도록 역할과 의사결정 절차를 명확히 합니다.
- **점검 및 조정:** 데이터와 성과 지표를 기반으로 AI 적용 현황을 점검하고, 변화에 맞춰 로드맵을 지속적으로 보완합니다.

변화의 신호

- **통합 로드맵 정립:** 명확한 AI 전략과 단계별 실행 계획이 수립됩니다.
- **공통 표준 운영:** 신규 AI 프로젝트가 공통 데이터, 보안, 리스크 기준에 따라 운영됩니다.
- **임원 차원의 검토:** 경영진이 AI 성과와 리스크를 정기적으로 점검합니다.
- **업무 구조 변화:** 개별 파일럿을 넘어 핵심 업무와 시스템에 AI가 내재화됩니다.

AI 리더의 사고방식

AI 확장은 단일 기술 프로젝트가 아니라 조직 전체의 재설계 과제입니다. 뛰어난 AI 모델이나 리더십 의지만으로는 충분하지 않습니다. 전략, 데이터, 기술, 거버넌스, 인력이 같은 방향으로 연결될 때 AI는 조직의 운영 방식에 내재화됩니다.

AI 리더는 각 사업 부문을 전사 전략과 연결하고, 공통 표준을 마련하며, 성과와 리스크를 함께 관리해야 합니다. AI를 단순한 도구가 아니라 조직 전반의 경쟁력을 높이는 운영 체계로 전환하는 것이 핵심입니다.

따라서 중요한 질문은 “어디에 AI를 적용할 것인가”에 그치지 않습니다. “AI를 전제로 조직의 일하는 방식과 의사결정 구조를 어떻게 다시 설계할 것인가”가 AI 전환 로드맵의 핵심 질문입니다.

제로 베이스(Zero-basing)는 기존의 관행과 경험에서 벗어나, 모든 것을 새롭게 바라보고 접근하는 혁신적 접근 방식입니다.

가치 측정은 AI의 효과를 사업 관점에서 확인하고, 확장 여부를 데이터에 근거해 판단하기 위한 체계입니다. 수익, 비용, 리스크, 고객 경험에 미치는 영향을 측정할 수 있어야 AI 확장도 가능합니다.

AI 가치는 명확한 기준 없이는 쉽게 드러나지 않습니다. AI 기능이 업무에 처음 적용될 때부터 사업 KPI, 기준값, 비교 대상, 비용 구조를 함께 설정해야 합니다. 그래야 AI가 실제로 매출을 늘렸는지, 비용을 줄였는지, 리스크 판단을 개선했는지 확인할 수 있습니다.

가치 측정은 단순한 모니터링이 아닙니다. AI가 업무 현장에서 어떤 성과를 만들고 있는지, 어떤 영역에서 개선이 필요한지, 어디에 추가 투자를 해야 하는지 판단하기 위한 운영 체계입니다. 성과 지표와 비용 지표, 위험 신호가 함께 관리될 때 AI 확장은 감이 아니라 근거에 기반해 이뤄질 수 있습니다.

특히 AI가 핵심 업무와 의사결정 과정에 들어올수록 측정 체계의 중요성은 커집니다. 단순히 사용량이나 처리 건수를 보는 것만으로는 충분하지 않습니다. AI가 수익성, 생산성, 품질, 고객 만족, 리스크 관리에 어떤 영향을 주는지 함께 확인해야 합니다. 이를 통해 기업은 성과가 낮은 과제는 조정하고, 성과가 입증된 과제는 더 빠르게 확산할 수 있습니다.

2026년에는 AI의 성과와 위험을 실시간에 가깝게 파악하는 역량이 중요해집니다. AI 에이전트와 자동화 기능이 확산되면서 기업은 수익과 비용, 품질, 보안, 규제 리스크를 지속적으로 점검해야 합니다. 이러한 측정 체계는 AI를 안정적으로 운영하고, 경영진이 확장 여부를 판단하는 핵심 근거가 됩니다.

왜 지금 중요한가

최근 기업들은 AI를 다양한 업무에 도입하고 있지만, 실제 사업 성과와 연결해 설명하는 데 어려움을 겪고 있습니다. 어떤 AI 기능이 수익을 늘리고, 어떤 기능이 비용을 줄이며, 어떤 기능이 리스크를 낮추는지 명확히 보여주지 못하면 투자는 지속되기 어렵습니다.

앞으로 AI 확장은 단순한 기술 도입이 아니라 투자 판단의 문제가 됩니다. 경영진은 AI가 만들어내는 성과와 비용, 위험을 함께 확인해야 하며, 이를 바탕으로 확장·중단·재투자 여부를 결정해야 합니다. 명확한 측정 체계는 AI 투자에 대한 신뢰를 높이고, 조직 내 확산 속도를 높이는 기반이 됩니다.

2026년에 변화하는 것

AI 가치 측정은 개별 프로젝트 평가를 넘어 시스템 수준의 운영 관리로 확대되고 있습니다:

- **모델 지표에서 사업 지표로:** 정확도, 처리 속도 같은 기술 지표를 넘어 매출, 비용, 고객 만족, 리스크 감소 등 사업 성과와 연결된 지표가 평가의 중심이 됩니다.
- **사용량을 넘어 가치 측정으로:** 토큰 사용량, 호출 횟수, 처리 건수만으로는 충분하지 않습니다. AI 사용이 실제 성과로 이어졌는지 확인하는 체계가 필요합니다.
- **실시간 운영 점검:** AI 에이전트와 자동화 기능이 늘어나면서 오류, 비용 증가, 이상 징후를 조기에 감지하는 운영 데이터가 판단의 기준이 됩니다.
- **확장 판단 기준 강화:** 핵심적인 소수의 가치·리스크 지표를 기준으로 AI 과제의 지속, 확대, 중단 여부를 결정합니다.

리더가 수익, 비용, 리스크를 명확히 파악할 수 있을 때 AI 확장은 더 빠르고 안정적으로 이뤄집니다

AI는 단순한 분석 도구를 넘어 업무를 직접 수행하는 시스템으로 확장되고 있습니다. 따라서 기업은 AI가 어디에서 성과를 만들고, 어디에서 비용과 위험을 발생시키는지 지속적으로 확인해야 합니다. 완전한 정보보다 중요한 것은 필요한 시점에 올바른 신호를 빠르게 포착하고 대응하는 능력입니다.

한국의 상황

국내 기업의 AI 도입은 빠르게 확산되고 있지만, AI가 실제로 어떤 가치를 만들어내는지 정량적으로 관리하는 체계는 아직 충분히 정착되지 않았습니다. 많은 AI 과제가 처리 시간 단축, 상담 품질 개선, 업무 자동화 등 개별 성과를 제시하고 있으나, 이를 매출 기여, 비용 절감, 리스크 감소, 고객 경험 개선과 연결해 지속적으로 추적하는 방식은 제한적입니다.

2026년 「AI 기본법」 시행과 공공·민간 AX 확산은 AI 성과와 리스크를 함께 측정해야 할 필요성을 높이고 있습니다. 앞으로 국내 기업은 AI 활용 현황을 단순히 도입 건수나 사용량으로 관리하는 데서 벗어나, 수익성, 운영 비용, 품질, 안전성, 규제 리스크를 함께 관리하는 성과 측정 체계를 갖춰야 합니다.

리더십 우선순위

- **가치 지표 연결:** AI 기능을 매출, 비용, 고객 만족, 리스크 등 사업 지표와 직접 연결하고 정기적으로 검토합니다.
- **실질적 성과 추적:** AI 적용 전후의 차이를 확인할 수 있도록 기준값, 비교 대상, 비용 구조를 함께 관리합니다.
- **위험 신호 관리:** 오류, 보안, 규제 리스크, 품질 저하 등 운영 과정에서 발생할 수 있는 위험 신호를 지속적으로 점검합니다.

변화의 신호

- **경제적 효과 추적:** 개별 AI 과제와 연결된 수익, 비용, 리스크 변화가 정기적으로 보고됩니다.
- **가시적인 단위 비용:** 컴퓨팅 비용, 데이터 비용, 운영 비용이 성과 지표와 함께 관리됩니다.
- **표준 증거 체계:** 모델 리스크, 개인정보 보호, 공정성, 품질 지표가 의사결정 자료로 활용됩니다.
- **업무 흐름 도입:** AI 사용 현황과 성과가 별도 보고서가 아니라 실제 업무 프로세스 안에서 관리됩니다.

AI 리더의 사고방식

AI 리더는 성과와 리스크를 동시에 측정하는 체계를 갖춰야 합니다. AI의 가치는 단순한 사용량이 아니라, 수익 개선, 비용 절감, 고객 경험 향상, 리스크 감소로 입증되어야 합니다.

확장 속도에 맞춰 무엇이 성과를 내고 있는지, 어디에서 비용과 위험이 발생하는지 지속적으로 확인해야 합니다. 이러한 측정 체계가 갖춰질 때 AI 확장은 더 안정적이고 신뢰 가능한 방식으로 이뤄질 수 있습니다.

업무 방식의 재설계

업무 방식의 재설계는 AI가 리더십, 업무 흐름, 직무, 역량과 문화 등 조직 운영의 핵심 요소를 어떻게 바꾸는지 살펴보는 관점입니다.

리더십 역량은 경영진이 AI의 가치와 리스크를 직접 이해하고 의사결정에 반영하는 능력을 다룹니다. 사람 중심 업무 설계는 AI를 기존 업무에 덧붙이는 것이 아니라, 사람이 더 가치 있는 판단과 창의적 역할에 집중할 수 있도록 업무 흐름을 다시 설계하는 과정입니다. 업무 혁신은 고정된 직무 중심 구조에서 벗어나 사람과 AI가 역할을 나누고 협업하는 방식으로 업무 구조가 전환되는 흐름을 설명합니다. 문화와 역량은 이러한 변화가 일회성 도입에 그치지 않고 조직의 일하는 방식으로 정착되도록 하는 기반이 됩니다.



AI가 경영 의사결정의 핵심 변수로 부상하면서 리더십의 기준도 달라지고 있습니다. 이제 리더에게 필요한 것은 AI 기술을 이해하는 수준을 넘어, AI가 만드는 가치와 리스크를 판단하고 전략과 운영에 반영하는 역량입니다.

과거에는 많은 기업이 AI를 기술 조직이 담당하는 전문 영역으로 여겼습니다. 그러나 2026년 현재 AI는 투자, 수익성, 리스크, 고객 경험, 운영 효율과 직접 연결되며 최고경영진의 의사결정 영역으로 들어오고 있습니다.

리더십 역량은 AI 모델의 세부 구조를 아는 데 그치지 않습니다. AI가 자사 사업에서 어떤 가치를 만들 수 있는지, 어디에 투자해야 하는지, 어떤 위험을 관리해야 하는지를 판단하는 능력입니다. 리더는 AI를 통해 고객 경험, 운영 효율, 혁신 속도가 어떻게 달라질 수 있는지 이해하고, 이를 전략적 선택으로 연결해야 합니다.

AI가 조직 운영과 의사결정 구조를 바꾸는 상황에서는 개별 프로젝트의 성공만으로 충분하지 않습니다. 리더는 AI 활용 방향과 기준을 제시하고, 투자 우선순위와 성과 기준, 리스크 관리 방식을 명확히 해야 합니다. AI가 단순한 실험에 머물지 않고 사업 성과로 이어지려면 리더의 판단과 책임이 함께 작동해야 합니다.

결국 리더십 역량은 “AI를 얼마나 알고 있는가”가 아니라, “AI를 어떻게 경영 판단과 실행으로 연결하는가”의 문제입니다. AI를 사업 전략의 중심에 두고, 성과와 리스크를 함께 관리하는 능력이 리더십의 중요한 기준이 되고 있습니다.

왜 지금 중요한가

2026년 현재 CEO와 경영진은 AI에 대한 실질적인 이해와 판단 능력을 요구받고 있습니다. 투자자, 규제기관, 고객은 기업이 AI를 얼마나 책임 있게 활용하고, 이를 얼마나 효과적으로 사업 성과와 연결하는지 주목하고 있습니다.

따라서 중요한 것은 형식적인 이해가 아니라 실행과 책임입니다. AI가 수익, 비용, 리스크, 고객 신뢰에 미치는 영향을 설명하고, 필요한 의사결정을 내릴 수 있어야 합니다.

2026년에 변화하는 것

AI 리더십은 성과와 리스크를 함께 판단하는 방향으로 변화하고 있습니다.

- **성과 중심의 AI 이해:** AI 논의가 모델 성능 중심에서 수익, 비용, 리스크, 고객 경험에 미치는 영향 중심으로 이동합니다.
- **전략 기반의 모델 선택:** 기술 파트너와 모델 선택은 단순한 성능 비교가 아니라 정책, 경제성, 리스크 기준에 따라 결정됩니다.
- **정기적 의사결정 체계:** 산발적 보고가 아니라, AI 성과와 리스크를 정기적으로 검토하고 확장·중단 기준을 명확히 합니다.
- **포트폴리오 단위 투자:** 개별 프로젝트가 아니라 단위 경제성, 회수 가능성, 전략적 중요도를 기준으로 AI 투자 포트폴리오를 관리합니다.
- **설계 단계의 신뢰 내재화:** 출시 기준과 통제 장치를 초기 설계 단계부터 반영해 속도와 신뢰를 함께 확보합니다.
- **사전 검증 강화:** AI 확장 전에 시뮬레이션과 테스트를 통해 주요 의사결정과 예외 상황을 점검합니다.

AI 문해력은 기술을 이해하는 능력입니다.

리더십 역량은 AI가 어디에서 가치를 만들고, 어떤 리스크를 만들며, 언제 개입해야 하는지를 판단하는 능력입니다.

한국의 상황

국내 기업에서 AI 리더십 역량은 모델의 기술적 작동 원리를 아는 것보다, AI가 실제 사업 성과와 운영 품질, 리스크에 어떤 영향을 주는지 읽어내는 능력에 가깝습니다. 특히 금융, 제조, 유통, 통신처럼 데이터와 운영 프로세스가 복잡한 산업에서는 AI 성과를 사업 지표와 연결해 판단하는 역량이 중요해지고 있습니다.

PwC 제28차 글로벌 CEO 설문에 따르면 생성형 AI에 대한 관심은 빠르게 높아졌지만, 많은 기업은 아직 AI 활용을 개별 사례나 기술 도입 수준에서 설명하는 데 머무르고 있습니다. 앞으로는 AI가 매출, 비용, 고객 경험, 운영 효율, 리스크 관리에 어떤 변화를 만드는지 경영진이 직접 이해하고 판단할 수 있어야 합니다.

리더십 우선순위

- **AI를 상시 경영 안건으로 관리:** AI 전략과 리스크를 정기 회의의 주요 안건으로 다루고, 투자와 성과를 함께 점검해야 합니다.
- **경영진의 AI 이해도 강화:** 리더가 근거 있는 의사결정을 내릴 수 있도록 AI의 사업적 영향, 리스크, 운영 기준에 대한 이해를 높여야 합니다.
- **성과 기반 보상 체계:** AI를 통해 창출된 실제 사업 성과가 경영진과 부문 리더의 평가 지표에 반영되도록 해야 합니다.
- **실무 참여 확대:** CEO와 사업부 리더는 주요 AI 프로젝트를 직접 후원하고, 현장 적용 과정에 참여해 AI가 조직의 핵심 우선순위를 분명히 해야 합니다.

변화의 신호

- **경영진 차원의 AI 논의:** AI가 도구 실험이 아니라 전략, 투자, 리스크 논의의 핵심 안건으로 다뤄집니다.
- **구체적 성과 언급:** CEO와 리더십 팀이 실적 발표나 내부 커뮤니케이션에서 AI의 실제 성과와 근거를 제시합니다.
- **작동하는 거버넌스:** AI 통제, 운영, 리스크 관리 체계가 구축되고 실제 의사결정에 반영됩니다.
- **근거 기반 투자 결정:** 주요 AI 투자가 명확한 사업 근거와 리스크 평가를 바탕으로 이뤄집니다.

AI 리더의 사고방식

AI에 대한 이해는 더 이상 기술 조직만의 과제가 아닙니다. 리더는 AI를 사업 성과, 고객 경험, 운영 효율, 리스크 관리와 연결해 판단할 수 있어야 합니다.

좋은 리더는 AI를 맹목적으로 도입하지 않습니다. 필요한 질문을 던지고, 전문가와 소통하며, 실제 도구가 현장에서 어떻게 작동하는지 확인합니다. 중요한 것은 모든 기술을 직접 아는 것이 아니라, 충분한 근거 없이 중요한 결정을 내리지 않는 판단력입니다.

AI 의사결정의 결과에 책임을 지고, 성과와 실패에서 배운 내용을 조직에 다시 반영하는 것. 이것이 AI 리더십 역량의 핵심입니다.

사람 중심 업무 설계

07

목표는 AI를 기존 업무에 덧붙이는 것이 아니라, 사람의 역할을 재정의하면서 업무 전 과정을 다시 설계하는 것입니다.

AI 업무 흐름이란 사람이 AI와 협업하는 방식으로 업무를 재구성하는 것을 의미합니다. 초기에는 이메일 작성, 문서 요약, 코드 생성처럼 단일 업무를 보조하는 방식으로 활용되는 경우가 많았습니다. 그러나 AI 활용이 고도화될수록 중요한 것은 특정 업무를 자동화하는 데 그치지 않고, 사람과 AI가 어떤 역할을 나누고 어떤 기준으로 협업할지 정하는 것입니다.

반복적이고 구조화된 업무는 AI가 처리하고, 사람은 판단과 책임이 필요한 영역에 집중하는 방식이 확산되고 있습니다. 효과적인 업무 설계는 AI의 강점과 사람의 강점을 구분하는 데서 출발합니다. AI는 초안 작성, 정보 탐색, 분류, 요약, 패턴 분석을 빠르게 수행할 수 있습니다. 반면 사람은 맥락을 해석하고, 예외 상황을 판단하며, 고객과 조직의 신뢰가 필요한 결정을 내리는 역할을 맡아야 합니다.

이러한 변화는 단순한 생산성 향상을 넘어 업무 방식 자체의 재설계를 요구합니다. AI가 업무의 일부를 담당하게 되면, 기존 직무와 역할, 승인 절차, 책임 구조도 함께 바뀌어야 합니다. 따라서 기업은 AI를 어디에 적용할지 뿐 아니라, 사람이 언제 개입하고 어떤 결정을 책임질지까지 명확히 해야 합니다.

결국 사람 중심 업무 설계의 핵심은 AI로 업무를 줄이는 것이 아니라, 사람이 더 높은 가치의 일에 집중할 수 있도록 업무를 재편하는 것입니다. AI가 반복 업무를 맡을수록 사람은 고객 판단, 창의적 문제 해결, 관계 형성, 윤리적 책임 등 인간 고유의 역량이 필요한 영역에 더 집중하게 됩니다.

왜 지금 중요한가

2026년 현재 많은 기업이 이메일 작성, 코드 생성, 상담 응대, 보고서 작성 등 다양한 업무에 AI 도구를 활용하고 있습니다. 그러나 단순 도입만으로는 생산성 향상이 지속되기 어렵습니다. 업무 흐름과 역할 분담, 책임 구조가 함께 바뀌지 않으면 AI 활용은 개별 편의 기능에 머물 가능성이 큼니다.

기업이 실질적인 생산성 향상을 얻으려면 업무를 사람과 AI의 협업 구조로 다시 설계해야 합니다. 동시에 직원들이 AI 협업을 신뢰하고, 자신의 전문성이 어디에서 더 큰 가치를 만드는지 이해할 수 있도록 변화 관리가 필요합니다. 이는 조직의 AI 성숙도를 높이는 핵심 조건입니다.

2026년에 변화하는 것

AI 업무 설계는 단순 자동화에서 사람-AI 협업 구조로 이동하고 있습니다.

- **전 과정 재설계:** 특정 업무에 AI를 적용하는 단계를 넘어, 사람과 AI가 함께 수행하는 업무 흐름을 처음부터 다시 설계합니다.
- **역할과 의사결정 기준 명확화:** AI가 초안 작성, 분석, 분류를 담당하고, 사람은 검토, 판단, 승인, 책임이 필요한 영역을 맡도록 역할을 구분합니다.
- **설계 단계의 신뢰 내재화:** 품질, 규정 준수, 검증 기준을 업무 설계 단계부터 반영해 AI 결과가 안정적으로 활용되도록 합니다.
- **AI 운영 방식 정착:** AI를 개별 파일럿이 아니라 업무 운영의 일부로 편입하고, 결과와 예외 상황을 지속적으로 관리합니다.

AI를 통해 업무 흐름을 다시 설계하고, 사람의 강점을 극대화하는 것이 핵심 목표입니다

한국의 상황

국내 기업은 생성형 AI와 자동화 도구를 빠르게 도입하고 있지만, 이를 사람의 역할과 업무 구조 재설계로 연결하는 단계는 아직 초기입니다. 많은 조직에서 AI는 문서 작성, 검색, 요약, 상담 지원 등 개별 업무 보조 수단으로 활용되고 있으나, 사람이 어떤 판단을 맡고 AI 결과를 어떻게 검토·승인할지에 대한 기준은 충분히 정립되지 않은 경우가 많습니다.

앞으로 중요한 것은 AI 활용을 단순한 업무 효율화가 아니라 사람과 AI의 역할을 새롭게 정의하는 과정으로 보는 것입니다. 국내 기업은 업무별로 AI가 맡을 일과 사람이 책임질 일을 구분하고, 검토·승인 절차, 오류 대응 기준, 직원 교육 체계를 함께 마련해야 합니다. 이를 통해 AI 활용은 일시적인 생산성 개선을 넘어, 사람의 판단과 전문성을 강화하는 업무 혁신으로 확장될 수 있습니다.

리더십 우선순위

- **대표 업무 선정:** 높은 비즈니스 가치를 가진 핵심 프로세스를 선정하고, 기존 흐름과 AI가 결합된 새로운 업무 방식을 설계합니다.
- **역할과 책임 구조 명확화:** 직원, 기술 전문가, 규제·보안 담당자가 함께 참여해 사람과 AI의 역할을 업무 처음부터 끝까지 정의합니다.
- **업무 기반 지원 체계 마련:** AI가 내재화된 업무 흐름을 지원하기 위해 필요한 데이터, 플랫폼, 도구를 제공하고, 실험과 개선이 가능한 운영 환경을 조성합니다.
- **조직 적응 지원:** 교육, 역할 재정의, AI 챔피언 육성 등을 통해 직원이 새로운 협업 방식에 안정적으로 적응할 수 있도록 지원합니다.

변화의 신호

- **처리 속도 개선:** 전체 업무 처리 시간이 의미 있게 단축됩니다.
- **업무 기준의 명문화:** AI가 포함된 업무 절차와 운영 가이드가 명확히 정리됩니다.
- **조직 문화 변화:** 직원이 AI를 단순한 도구가 아니라 함께 일하는 협업 대상으로 인식하고, 활용 과정에서 심리적 안정감을 갖게 됩니다.
- **개선 효과의 다면화:** 처리 속도뿐 아니라 품질, 고객 만족도, 직원 경험 등 여러 영역에서 협업을 통한 업무 개선 효과가 확인됩니다.

AI 리더의 사고방식

AI를 기존 프로세스에 추가하는 도구로만 봐서는 안 됩니다. AI를 전제로 업무 흐름, 역할, 책임 구조를 다시 설계할 때 실질적인 변화가 만들어집니다.

AI 리더는 어떤 업무를 자동화할지보다, 사람이 어디에서 더 큰 가치를 만들 수 있는지를 먼저 판단해야 합니다. 반복 업무는 AI가 맡고, 사람은 판단·관계·창의성·책임이 필요한 영역에 집중하도록 업무를 설계해야 합니다.

중요한 것은 AI가 사람을 대체하는 구조가 아니라, 사람의 역량을 확장하는 구조를 만드는 것입니다. 이 관점이 자리 잡을 때 AI는 조직의 생산성을 높이는 동시에, 사람 중심의 업무 방식을 강화하는 기반이 됩니다.

AI는 인간의 잠재력을 대체하는 것이 아니라 이를 증폭시키는 역할을 수행합니다.

이 과정에서 모든 AI 활용은 인간의 통제권을 전제로 설계되어야 하며, 이는 '인간이 주도하는 AI(Humans at the Helm)' 원칙과도 일치합니다.

**AI는 업무 속도 향상을 넘어, 업무 구조와 역할 분담을 바꾸고 있습니다.
사람과 AI가 각자의 강점을 발휘하도록 업무 전 과정을 다시 설계하는 것이 핵심입니다.**

업무의 성격은 도구가 아니라 구조의 차원에서 바뀌고 있습니다. AI가 실행과 의사결정 지원의 영역으로 들어오면서, 기존 직무는 단순히 보완되는 것이 아니라 새롭게 설계되고 있습니다. 업무는 더 빠르고 더 많이 처리되는 데 그치지 않고, 역할과 책임, 판단의 기준까지 다시 정의되고 있습니다.

과거 업무는 고정된 직무와 절차를 중심으로 운영되었습니다. 그러나 이제 사람과 AI가 업무를 나누고, 의사결정은 더 넓게 분산되며, 성과는 사람과 AI의 협업을 통해 만들어집니다. 사람은 판단, 관계 형성, 윤리적 책임과 같이 맥락이 필요한 영역에 강점을 갖고, AI는 속도, 정보 처리, 패턴 분석, 요약과 종합에서 강점을 보입니다.

따라서 기업은 AI를 기존 업무에 단순히 추가하는 데서 벗어나야 합니다. 중요한 것은 사람이 가장 큰 가치를 만들 수 있는 영역과 AI가 더 효과적으로 수행할 수 있는 영역을 구분하고, 이를 기준으로 업무 흐름과 역할 체계를 다시 설계하는 것입니다.

이 변화의 중심에는 파트너십이 있습니다. 사람과 AI의 협업 관계가 명확하지 않으면 병목과 재작업이 늘어날 수 있습니다. 반대로 협업 구조가 잘 설계되면 자율성, 속도, 신뢰성이 함께 높아집니다. 앞으로의 업무 혁신은 사람의 판단력과 맥락 이해, AI의 처리 능력과 기억력, 분석 역량을 결합하는 데서 출발합니다.

의사결정 구조도 함께 바뀌어야 합니다. 모든 AI 결과를 사람이 일일이 승인하는 방식으로는 확장에 한계가 있습니다. 리스크 수준에 따라 검토 기준을 다르게 두고, 근거와 통제 장치를 업무 안에 포함하며, 사람이 언제 개입해야 하는지를 명확히 해야 합니다. 신뢰는 역할이 얼마나 명확한지, 사람이 얼마나 자주 개입해야 하는지에 따라 달라집니다.

왜 지금 중요한가

AI 확산 속도가 이를 안정적으로 운영할 수 있는 업무 체계의 정비 속도를 앞서고 있습니다. 많은 조직이 기존 업무 구조는 그대로 둔 채 AI 도구를 추가하거나 일부 업무만 자동화하고 있습니다. 그 결과 산출량은 늘어나지만 품질 편차와 검토 부담은 커지고, AI가 만든 결과를 사람이 사후에 확인하고 수정하는 구조가 반복되고 있습니다.

AI는 단순히 처리 속도만 높이는 것이 아닙니다. 의사결정, 예외 상황, 검토 업무까지 함께 늘릴 수 있습니다. 따라서 지금 중요한 것은 AI 도구를 더 많이 도입하는 것이 아니라, 사람과 AI가 어떤 역할을 나누고 어떤 기준으로 협업할지 업무 구조를 다시 설계하는 것입니다. 이 설계에 따라 AI는 조직 역량을 높이는 기반이 될 수도, 업무 복잡성을 키우는 요인이 될 수도 있습니다.

2026년에 변화하는 것

업무 혁신은 단순 자동화에서 사람-AI 협업 구조로 이동하고 있습니다:

- **협업 구조 우선 설계:** AI 도입에 앞서 사람과 AI가 맡을 역할, 사람의 판단이 필요한 지점, 검토와 승인 기준을 먼저 정합니다.
- **일상 업무 속 AI 활용 정착:** AI는 초안 작성, 정보 정리, 분류, 실행 지원을 맡고, 사람은 검토와 판단, 예외 상황 대응에 집중합니다.
- **관리자 역할의 전환:** 관리자는 사람과 AI가 함께 일하는 과정을 조율하고, 책임 범위와 의사결정 기준을 명확히 하는 역할을 맡습니다.
- **고부가가치 업무로 전환:** AI로 반복 업무 부담이 줄어든 만큼, 고객 대응, 문제 해결, 업무 개선 등 더 가치 있는 활동에 역량을 집중합니다.

AI는 고정된 직무 중심의 업무 방식에서 벗어나, 보다 유연한 역할 중심의 업무 구조를 만들어가고 있습니다.

수동적인 업무 절차는 사람-AI 협업 체계로 전환되고 있습니다.
이 흐름에 맞춰 업무를 설계하는 조직은 단순한 생산성 개선을 넘어, 더 지속 가능하고 의미 있는 업무 혁신을 만들어낼 수 있습니다.

한국의 상황

한국의 AI 역량 강화 전략은 정부 주도의 디지털 인재 양성 정책과 기업의 실무 중심 교육을 연계하는 방향으로 발전하고 있습니다.

국내 기업들은 AI를 단순한 생산성 도구가 아니라 개인의 경력과 함께 축적되는 '휴대 가능한 역량(Portable Skills)'으로 인식하며 인적 자본의 가치를 높이고 있습니다.

또한 노사 관계에서도 AI 도입에 따른 고용 불안을 완화하기 위해 기술 전환기 고용 보장과 직무 재배치를 핵심 의제로 다루고 있습니다. 이는 AI가 인간을 대체하는 것이 아니라 인간의 역량을 강화하는 생산성 향상 모델로 발전하고 있음을 보여줍니다.

리더십 우선순위

- **협업 관계 설계:** 사람과 AI가 함께 수행하는 업무에서 역할, 판단 기준, 검토 절차, 책임 범위를 명확히 정의합니다.
- **협업팀 역량 강화:** 관리자가 사람-AI 협업을 이끌 수 있도록 촉진, 성과 관리, 신뢰 형성 역량을 함께 강화합니다.
- **공유 시스템 구축:** 사람과 AI가 같은 업무 맥락에서 협업할 수 있도록 데이터, 도구, 업무 기준을 통합해 제공합니다.
- **협업 효과 측정:** 생산성뿐 아니라 품질, 관계의 안정성, 직원 경험, 고객 만족도까지 함께 측정합니다.

변화의 신호

- **역할 구분:** 구성원은 자신이 직접 판단할 업무, AI에 맡길 업무, 함께 처리할 업무를 명확히 이해합니다.
- **실제 업무 적용:** AI 활용이 별도 실험이나 개인 선택에 머물지 않고, 문서 작성·분류·검토 등 실제 업무 절차에 포함됩니다.
- **개입 기준 마련:** AI는 정해진 기준에 따라 반복 업무를 처리하고, 사람은 예외 상황이나 중요한 판단이 필요한 사안에 개입합니다.
- **현장 피드백 반영:** 팀별 활용 경험이 공유되고, 이를 바탕으로 업무 절차와 AI 사용 기준이 지속적으로 보완됩니다.

AI 리더의 사고방식

AI 리더는 사람만 관리하는 것이 아니라, 사람과 AI의 협업 관계를 이끌어야 합니다. 리더십의 초점은 인적 자원 관리에서 사람과 AI가 함께 더 큰 역량을 발휘하도록 업무 관계를 조율하는 방향으로 이동하고 있습니다.

중요한 것은 기술을 먼저 배포하는 것이 아니라, 사람과 AI가 어떤 역할을 맡고 어떻게 상호 보완할지를 먼저 설계하는 것입니다. 업무 혁신의 성패는 사람과 AI가 각자의 강점을 최대한 발휘할 수 있는 협업 구조를 만드는 데 달려 있습니다.

이제는 AI와 함께 성과를 낼 수 있는 환경과 역량을 구축해야 할 때입니다.

AI 경제에서 일자리를 지키는 가장 확실한 방법은 자동화 가능한 업무를 넘어 더 높은 가치를 만들어내는 것입니다.

PwC 글로벌 AI 일자리 바로미터에 따르면, AI 도입이 활발한 산업일수록 생산성 향상이 빠르게 나타나고 있으며, AI 활용 역량을 갖춘 인재는 임금 측면에서도 프리미엄을 얻고 있습니다.

그러나 개인의 역량만으로는 충분하지 않습니다. AI 성숙도는 구성원의 기술 수준뿐 아니라, 이를 실제 업무에 적용할 수 있는 조직 환경에 의해 좌우됩니다. 구성원이 AI를 활용해 실험하고, 결과를 검토하며, 시행착오를 통해 일하는 방식을 바꿀 수 있어야 AI가 실질적인 성과로 이어질 수 있습니다.

핵심은 AI를 단순한 도구가 아니라 함께 일하는 업무 파트너로 활용하는 능력입니다. 이를 위해서는 프롬프트 설계, 결과 검토, 근거 확인, 오류 식별, 사람의 개입이 필요한 상황을 판단하는 역량이 필요합니다. 이러한 역량은 조직이 활용을 장려하고 학습을 지원할 때 비로소 성과로 연결됩니다.

성과를 내는 AI 조직 문화는 자신감에서 출발합니다. 구성원들이 AI를 위협이 아니라 자신의 역량을 확장하는 수단으로 받아들일 때 업무 방식의 변화가 시작됩니다. 조직은 윤리적 기준을 유지하면서도 실험을 장려하고, 검증된 활용 방식을 실제 업무로 확산할 수 있는 환경을 마련해야 합니다.

결국 문화와 역량은 하나의 과제입니다. 역량이 있어도 문화가 뒷받침되지 않으면 변화는 확산되지 않습니다. 반대로 문화적 준비 없이 AI 도입만 강화하면 업무 혼선과 피로감이 커질 수 있습니다. 구성원이 자신의 성장과 조직의 성과가 연결되어 있다고 느낄 때, AI 활용은 일상적인 행동 변화로 이어집니다.

왜 지금 중요한가

AI 확산의 다음 과제는 새로운 도구의 도입이 아니라, AI와 함께 일하는 방식을 조직 전반에 정착시키는 것입니다.

여기에는 AI 리터러시, 프롬프트 설계, 결과 검토, 비판적 판단, 오류 탐지, 사람의 개입 기준 등이 포함됩니다. AI에 무엇을 맡기고, 언제 사람이 판단해야 하는지를 구분하는 능력이 중요해지고 있습니다.

자동화가 확산될수록 인간 고유의 역량도 다시 주목받고 있습니다. 비판적 사고, 문제 해결 능력, 맥락 판단, 커뮤니케이션, 가치 판단은 AI가 대체하기 어려운 영역입니다. 앞으로의 경쟁력은 AI 활용량이 아니라, 사람과 AI의 협업 방식과 이를 통한 성과 창출 역량에 달려 있습니다.

2026년에 변화하는 것

조직이 AI를 본격적으로 확장하기 시작하면 문화와 역량은 더 이상 부가적인 요소가 아닙니다:

- **문화적 준비:** 파일럿을 실제 업무로 확장하기 위해 공통의 가치 기준과 의사결정 방식을 정비합니다.
- **학습의 일상화:** 정규 교육만으로는 변화 속도를 따라가기 어려워지면서, 현장 중심의 코칭과 실습 중요성이 커집니다.
- **AI 협업 중심의 역할 정립:** AI 활용 역량은 직무 수행의 핵심 요소가 되며, 역할 정의와 성과 관리에도 반영됩니다.
- **변화 관리의 전환:** AI 사용을 독려하는 데 그치지 않고, 새로운 업무 방식에 맞춰 역할과 책임, 협업 기준을 재정리합니다.

AI가 차량이라면, 역량은 운전면허와 같습니다.

운전면허가 차량을 안전하게 다루기 위한 기본 조건이듯,
AI 활용에도 기본 역량이 필요합니다.
구성원은 AI와 협업하는 방법을 이해하고, AI의 결과를
검토·수정·거부할 수 있어야 합니다. 이러한 역량이 확산될수록
업무 속도와 신뢰도는 높아지지만, 최종 책임은 여전히 사람에게 남습니다.

한국의 상황

한국은 빠른 기술 수용력을 갖추고 있지만, 위계적 조직문화로 인해 AI 도입과 실제 업무 혁신 사이에 간극이 생길 수 있습니다. 사람과 AI가 함께 일하는 조직으로 전환하려면 시행착오와 학습 경험을 공유할 수 있는 문화가 필요합니다.

AI를 단순한 자동화 수단으로만 볼 것이 아니라, 확보된 시간을 더 높은 수준의 판단과 기획에 활용하도록 업무 구조와 평가 방식을 함께 바꿔야 합니다. 사람의 개입 지점과 책임 소재를 명확히 하는 것도 중요합니다.

리더십 우선순위

- **협업 원칙 정립:** 사람이 반드시 판단해야 하는 영역과 AI가 수행할 수 있는 영역을 명확히 구분해야 합니다. 역할 정의, 의사결정 권한, 책임 기준도 함께 정리해야 합니다.
- **AI 활용 역량 강화:** 관리자는 구성원이 AI를 실제 업무에 효과적으로 활용할 수 있도록 학습과 실습 기회를 제공해야 합니다. 단순한 도구 교육을 넘어, 업무 맥락에서 AI 결과를 검토하고 활용하는 능력을 키우는 것이 중요합니다.
- **공유 시스템 구축:** 사람과 AI가 원활하게 협업할 수 있도록 공통의 업무 기준과 협업 방식을 마련해야 합니다. 정보가 단절되지 않고, 구성원이 같은 기준으로 AI를 활용할 수 있는 환경이 필요합니다.

변화의 신호

- **명확한 역할 구분:** 구성원들은 자신이 담당하는 업무와 AI가 담당하는 업무를 명확하게 이해하고 있으며, 인간과 AI의 역할 경계가 조직 전반에 걸쳐 공유됩니다.
- **원활한 협업:** 인간-AI는 정리된 업무와 공유 시스템을 기반으로 협업하며, 업무 인계와 의사결정 과정이 자연스럽게 연결됩니다.
- **신뢰 기반의 자율성:** AI는 정해진 기준과 범위 안에서 자율적으로 업무를 수행하고, 인간은 지속적인 감시와 검토 대신 예외 상황과 고부가가치 의사결정에 집중합니다.

AI 리더의 사고방식

리더의 역할은 더 이상 사람만 관리하는 데 머물지 않습니다. 이제는 인간과 AI의 협업을 조율하고, 어느 한쪽만으로는 달성할 수 없는 역량을 끌어내는 것이 중요합니다.

이를 위해서는 기술 도입보다 협업 구조를 먼저 설계해야 합니다. 성공적인 조직은 사람과 AI가 각자의 강점을 가장 효과적으로 발휘할 수 있는 업무 환경을 구축합니다. 결국 경쟁력의 원천은 기술 자체가 아니라, 사람과 AI가 함께 성과를 만들어내는 협업의 질에 달려 있습니다.

지능형 시스템 구축

AI를 현실로 구현하는 기술 아키텍처

AI를 실제 업무에서 작동하게 하려면 모델, 데이터, 에이전트, 운영 체계, 컴퓨팅 자원이 하나의 구조로 연결되어야 합니다.

파운데이션 모델은 AI 활용의 기본 역량을 제공하고, 에이전트형 AI 시스템은 업무 실행 범위를 넓힙니다. 맥락은 기업의 데이터와 정책을 AI 판단에 연결하며, 시뮬레이션 환경은 적용 전 다양한 상황을 검증합니다.

AI 운영(AI Ops)은 모델과 에이전트가 안정적으로 배포 · 관리되도록 지원하고, 컴퓨팅 자원 관리는 비용과 성능, 용량을 함께 고려해 지속 가능한 운영 기반을 마련합니다.



파운데이션 모델

파운데이션 모델은 주요 AI 활용 사례의 기반이 되는 범용 AI 모델입니다.

다만 모델 접근성만으로는 더 이상 차별화가 어렵습니다. 앞으로의 경쟁력은 모델 성능보다 경제성, 신뢰성, 통제 가능성을 갖춘 운영 역량에서 결정됩니다.

파운데이션 모델은 글쓰기, 코딩, 분석, 고객 응대 등 다양한 AI 활용 사례의 기반이 되어 왔습니다. 초기 경쟁은 모델 규모와 성능을 중심으로 전개됐습니다. 그러나 주요 모델 간 성능 격차가 줄고 선택지가 늘어나면서, 기업의 관심은 특정 모델을 보유하는 것에서 업무에 맞는 모델을 선별하고 조합하는 방향으로 이동하고 있습니다.

이제 중요한 것은 하나의 고성능 모델에 모든 업무를 맡기는 것이 아닙니다. 단순한 요약이나 분류 작업은 비용 효율이 높은 모델로 처리하고, 복잡한 분석이나 중요한 의사결정에는 더 높은 성능의 모델과 검토 절차를 적용하는 방식이 확산되고 있습니다. 모델 자체의 성능보다 업무 목적, 비용, 속도, 리스크에 맞춰 모델을 배분하고 운영하는 능력이 중요해지고 있습니다.

파운데이션 모델의 가치는 기업의 데이터, 업무 맥락, 보안 기준, 운영 체계와 결합될 때 실질적인 성과로 이어집니다. 같은 모델을 사용하더라도 어떤 데이터를 연결하는지, 어떤 기준으로 답변을 검토하는지, 민감한 정보가 어디에서 처리되는지에 따라 결과의 품질과 신뢰도는 달라집니다. 모델 성능이 높더라도 부정확한 데이터, 불명확한 책임 구조, 취약한 통제 기준 위에서 운영된다면 기대한 효과를 내기 어렵습니다.

따라서 기업은 파운데이션 모델을 단일 기술 자산이 아니라 운영 포트폴리오로 관리해야 합니다. 업무별로 적합한 모델을 선택하고, 필요할 때 다른 모델로 전환할 수 있는 구조를 갖춰야 합니다. 동시에 사용 비용, 처리 속도, 오류율, 재작업 수준, 보안 요건을 함께 관리해야 합니다. 파운데이션 모델을 안정적으로 활용하려면 모델 선택 기준뿐 아니라, 이를 운영하고 통제하는 체계가 함께 마련되어야 합니다.

왜 지금 중요한가

파운데이션 모델이 보편화될수록 기업의 과제는 모델 성능 비교보다 운영 기준 정립에 가까워지고 있습니다. 업무 특성, 데이터 민감도, 비용 구조, 리스크 수준에 맞춰 모델을 구분하지 않으면 비용 변동성, 공급자 종속, 품질 편차, 보안 리스크가 커질 수 있습니다.

기업은 모델을 선택할 때 성능만이 아니라 업무 중요도, 데이터 민감도, 비용 구조, 규제 요건, 사람의 검토 필요성을 함께 고려해야 합니다. 파운데이션 모델이 보편화될수록 경쟁력은 모델 접근성보다 모델을 선택하고 운영하는 기준에서 만들어 집니다.

2026년에 변화하는 것

파운데이션 모델은 컴퓨팅 자원처럼 보편화되고, 경쟁의 초점은 모델 성능 자체보다 경제성, 신뢰성, 통제 가능성으로 이동합니다.

- **작업별 모델 배분:** 단순 업무는 비용 효율이 높은 모델로 처리하고, 복잡한 분석이나 중요한 판단에는 더 높은 성능의 모델과 검토 절차를 적용합니다.
- **모델 사용 기준 정립:** 업무 유형, 데이터 민감도, 리스크 수준에 따라 사용할 수 있는 모델과 검토·승인 기준을 정합니다.
- **데이터 내부 처리 확대:** 프라이빗 추론, 엣지 배포, 국내 리전 활용이 확대되며, 민감 데이터는 가능한 한 데이터가 있는 위치에서 처리하는 방향으로 전환됩니다.
- **소형·특화 모델 활용 확대:** 특정 업무에서 충분한 성능을 보이는 소형 모델과 특화 모델의 활용이 늘어나고, 효율성이 중요한 경쟁 기준이 됩니다.

성공은 가장 강력한 파운데이션 모델에 접근하는 것만으로 이루어지지 않습니다.

업무에 맞는 모델을 선택하고, 신뢰할 수 있는 데이터와 사람의 검토를 결합하며, 비용과 성과를 함께 관리하는 체계를 갖추 때 경쟁력이 생깁니다. 중요한 것은 모델 성능보다 맥락, 데이터 관리, 운영 기준, 비용 대비 효과입니다.

한국의 상황

국내에서는 독자 파운데이션 모델과 소버린 AI 논의가 활발해지면서, 기업의 모델 선택 기준도 더 중요해지고 있습니다. 한국어 성능, 산업 데이터 활용, 보안 요건, 규제 대응, 비용 구조를 함께 고려해야 하기 때문입니다.

국내 기업은 성능이 높은 모델을 도입하는 데 그치지 않고, 업무 특성에 맞는 모델을 선별해 활용해야 합니다. 고객 정보, 금융·의료 데이터, 내부 문서처럼 민감한 정보를 다루는 업무에서는 데이터가 어디에서 처리되는지, 어떤 권한과 통제 기준이 적용되는지가 핵심 기준이 됩니다.

앞으로 국내 기업의 과제는 가장 강력한 모델을 찾는 것이 아니라, 업무 목적과 리스크 수준에 맞게 모델을 조합하고 운영하는 것입니다. 프라이빗 추론, 국내 리전 활용, 엣지 배포, 소형·특화 모델 적용은 비용과 보안, 성능을 함께 관리하기 위한 중요한 선택지가 될 것입니다.

리더십 우선순위

- **모델 전환 가능성 확보:** 특정 모델이나 공급자에 과도하게 의존하지 않도록, 업무에 맞는 모델을 선택하고 필요할 때 전환할 수 있는 구조를 마련해야 합니다.
- **검증 체계 정비:** 모델 도입 전 테스트, 취약점 점검, 운영 기준, 사고 대응 절차를 마련해 실제 업무 적용에 따른 리스크를 줄여야 합니다.
- **비용과 성과 관리:** 모델별 사용 비용, 처리 시간, 품질, 오류율을 함께 확인해 비용 대비 효과를 지속적으로 점검해야 합니다.

변화의 신호

- **정책 기반 모델 선택:** 개인 선호나 관행이 아니라, 업무 기준과 리스크 수준에 따라 모델이 선택됩니다.
- **소형·특화 모델 활용 확대:** 반복적이고 범위가 명확한 업무에는 대형 모델보다 비용 효율이 높은 모델이 적용됩니다.
- **고위험 업무 검토 강화:** 민감한 정보나 중요한 판단이 포함된 업무에는 사람의 검토와 승인 절차가 함께 적용됩니다.
- **업무별 비용·품질 관리:** 리더가 작업별 비용, 품질, 오류율, 재작업 수준을 확인하고 모델 활용 방식을 조정합니다.

AI 리더의 사고방식

경영진은 파운데이션 모델을 빠르게 교체되는 기술 상품으로만 보아서는 안 됩니다. 핵심은 개별 모델의 성능보다, 모델을 안정적으로 운영하고 통제할 수 있는 체계를 갖추는 데 있습니다.

신뢰성, 경제성, 거버넌스는 모델 자체에서 자동으로 만들어지지 않습니다. 업무 목적에 맞는 모델을 선택하고, 데이터·비용·보안·검토 기준을 함께 관리할 때 실제 성과로 이어집니다. 따라서 AI 리더십의 초점은 최고 성능 모델의 확보가 아니라, 조직의 업무에 맞는 모델을 안전하고 효율적으로 운영하는 체계 구축에 놓여야 합니다.

에이전트형 AI 시스템

에이전트형 AI는 여러 단계의 작업을 스스로 계획하고 실행하는 AI 시스템입니다. 단순히 답을 제안하는 데 그치지 않고, 도구와 데이터를 활용해 실제 업무 수행까지 지원합니다.

2025년 한 해 동안 AI 분야에서 가장 주목받은 흐름 중 하나는 단일 목적 에이전트에서 여러 에이전트가 함께 작동하는 체계로의 전환이었습니다. 초기 에이전트는 이메일 처리, 고객 응대, 데이터 조회처럼 제한된 업무에서 가능성을 보였습니다. 그러나 권한 설정, 시스템 연동, 예외 처리, 책임 기준이 충분히 갖춰지지 않은 상태에서는 전사적으로 확장되기 어려웠습니다.

앞으로의 과제는 에이전트를 신뢰할 수 있는 운영 체계 안에서 관리하는 것입니다. 에이전트는 더 이상 하나의 챗봇이나 단일 모델이 아니라, 각자 맡은 역할에 따라 도구를 호출하고 정보를 주고받으며 필요한 경우 사람에게 넘기는 실행 단위로 발전하고 있습니다. 따라서 에이전트형 AI 시스템은 기술 도입 과제가 아니라 운영 구조와 책임 체계를 함께 설계해야 하는 과제입니다.

이 변화가 중요한 이유는 기존 자동화 방식이 한계에 이르고 있기 때문입니다. 과거의 자동화는 정해진 규칙과 절차를 기반으로 반복 업무를 처리하는 데 강점이 있었습니다. 반면 에이전트형 AI는 정보를 찾고, 업무 맥락을 파악하며, 필요한 도구를 호출하고, 사람의 판단이 필요한 지점을 구분할 수 있습니다. 이를 통해 단편적인 자동화를 넘어 여러 업무 흐름을 연결하는 실행 체계로 확장될 수 있습니다.

에이전트는 사람의 업무 방식을 바꾸는 동시에, 관리자의 역할도 바꿉니다. 사람은 정책 판단, 예외 처리, 최종 검토처럼 책임이 필요한 영역에 집중하고, 에이전트는 반복적인 확인, 상태 업데이트, 자료 정리, 절차 실행을 담당할 수 있습니다. 앞으로 인력 운영은 사람의 수를 계획하는 데 그치지 않고, 사람과 에이전트가 어떤 역할을 나누고 어떻게 협업할지 설계하는 방향으로 확대될 것입니다.

이때 중요한 것은 에이전트에 맡길 업무와 사람이 직접 판단해야 할 업무를 명확히 구분하는 것입니다. 역할, 권한, 책임 범위, 인수 인계 기준, 중단 조건이 정리되어야 에이전트가 독립적으로 일하더라도 조직의 통제 범위 안에서 안정적으로 운영될 수 있습니다.

2026년에 변화하는 것

단기적으로 성과가 잘 드러나는 단일 에이전트 사례도 계속 나오겠지만, 핵심은 여러 에이전트를 안정적으로 조율하고 운영하는 체계로 이동하고 있습니다.

- **복수 에이전트 체계:** 단일 에이전트가 모든 일을 처리하는 방식에서 벗어나, 여러 에이전트가 역할을 나누고 서로 연계되는 구조로 발전합니다.
- **상호운용성 기준 확대:** 에이전트가 여러 업무 시스템과 연결되면서, 에이전트 간 연동 방식과 표준 인터페이스의 중요성이 커집니다.
- **권한 기반 거버넌스 강화:** 에이전트의 시스템 접근 권한, 인증, 기록 관리가 중요해집니다. 누가, 언제, 어떤 도구를 호출했는지 추적할 수 있어야 합니다.
- **운영 중 통제 체계 정립:** 에이전트가 실행 중인 업무 상태와 의사결정 과정을 실시간으로 확인하고, 필요할 때 즉시 개입할 수 있는 기준이 마련됩니다.
- **현장 실행 단계의 조율:** 에이전트가 실제 고객 접점과 업무 현장에서 작동하면서, 맥락 이해와 예외 처리, 사람으로의 인계가 성과를 좌우하게 됩니다.

2026년에는, 에이전트를 조직에 도입하는 과정이 실제 팀원을 맞이하는 과정과 비슷해 질 것입니다

권한 검증, 역할 정의, 업무 배치, 성과 미달 시 권한 회수와 역할 종료까지
에이전트 관리에는 실제 인력 관리와 유사한 절차가 필요합니다.

한국의 상황

국내 기업의 자동화는 단순한 규칙 반복에서, 스스로 인지 · 계획 · 실행하는 에이전트형 AI로 빠르게 이동하고 있습니다. 기업용 AI 에이전트 시장은 향후 수년간 폭발적으로 성장할 것으로 전망되며, 주요 기업과 금융기관은 청구 처리와 보고서 작성 같은 최종 업무에 에이전트를 배치해 실질적인 비용 절감과 생산성 향상을 거두기 시작했습니다.

이 전환의 핵심은 에이전트를 '디지털 워커'로 다루는 데 있습니다. 자동화 시스템 계정에 대한 권한 제어와 자율 작업 추적 이 핵심 보안 쟁점으로 떠오르면서 최소 권한과 감사 로그 중심 의 거버넌스가 정비되고 있습니다. 이에 따라 관리자는 사람과 디지털 인력을 통합해 감독하고 평가하는 혼합 인력 관리 역량을 요구받고 있습니다.

리더십 우선순위

인력 계획에 디지털 워커(AI)의 역량이 포함되면서 최고인사 책임자(CPO)의 역할은 더욱 확대됩니다. 역할 설계 단계에서는 명확한 위임 경계를 설정해야 하며, 관리자는 자신이 직접 수행하지 않은 업무를 감독하기 위한 루틴이 필요합니다.

조직적 확장에 도달할 준비가 되었을 때:

- **'디지털 워커' 관리 체계 구축:** 에이전트별 독립 식별자(ID), 역할 기반 권한, 명확한 소유권을 수립하여 모든 에이전트가 정의된 범위 내에서 감독과 책임을 갖도록 합니다.
- **하이브리드 팀 운영 기준 정비:** AI 에이전트의 업무 범위를 명확히 규정하여, 직원들이 '내가 할 일'과 'AI가 할 일'을 혼 선 없이 파악하도록 직무 기준과 관리 프로세스를 정비합니다.

변화의 신호

- **운영 범위 설정:** 에이전트가 사전 정의된 서비스 수준(SLA) 내에서 안정적으로 구동되도록 제어하며, 예외 상황 발생 시 구조화된 에스컬레이션(상위 보고 및 조치) 경로를 엄격히 준수합니다.
- **확장의 안정성:** 에이전트 도입 규모가 확대되더라도 시스템 오류, 재작업 리스크, 컴플라이언스 이슈의 증가 없이 전반 적인 운영 성과를 지속적으로 개선합니다.
- **보안 및 감사 추적 강화:** 상세 구동 로그와 독립된 에이전트 식별자(ID)를 기반으로, 보안·리스크 관리 팀이 침해 사고 발생 시 '행위자, 수행 작업, 일시, 목적(5W1H)'을 신속하게 추적 · 대응할 수 있도록 합니다.

AI 리더의 사고방식

AI 리더는 에이전트를 새로운 실행 주체로 인식해야 합니다. 이는 단순히 기술 기능을 배포하는 문제가 아니라, 업무 성과 와 운영 리스크를 함께 관리하는 문제입니다.

앞으로의 에이전트는 단순한 보조 도구가 아니라, 일부 업무를 직접 수행하는 디지털 실행 단위가 될 것입니다. 사람은 최종 판단과 책임이 필요한 영역에 집중하고, 에이전트는 정해진 권한과 기준 안에서 반복 업무와 절차 실행을 맡게 됩니다.

따라서 리더는 에이전트에도 사람과 마찬가지로 역할, 책임, 권한, 성과 기준을 부여해야 합니다. 목표 설정, 권한 관리, 모 니터링, 성과 점검, 복구 절차까지 포함한 운영 체계를 갖출 때 에이전트형 AI 시스템은 안정적인 생산성 기반으로 자리 잡을 수 있습니다.

맥락(Context)

맥락은 기업의 데이터, 문서, 정책을 AI와 연결해 조직의 기준에 맞는 답변과 판단을 가능하게 합니다. 맥락이 갖춰질 때 AI는 더 정확하고 신뢰할 수 있는 방식으로 업무를 수행할 수 있습니다.

파운데이션 모델은 공개된 지식에는 강점을 갖고 있지만, 개별 기업의 전략, 리스크 기준, 고객과의 약속, 실제 업무 과정에서 축적된 판단 기준까지 스스로 이해하지는 못합니다. 기업마다 사업 방식과 고객 요구, 규제 환경이 다르기 때문에 공개 지식만으로는 조직의 실제 맥락에 맞는 답변을 내기 어렵습니다.

따라서 AI 답변과 판단을 기업의 실제 상황에 맞게 조정하려면 기업 고유의 데이터, 문서, 정책, 의사결정 이력을 모델과 연결해야 합니다. 이때 AI 시스템에 제공되는 맥락은 모델이 조직의 업무 방식과 판단 기준을 이해하도록 돕는 핵심 정보 체계가 됩니다.

맥락은 단순히 데이터를 많이 넣는다고 만들어지지 않습니다. 어떤 정보가 최신이고 신뢰할 수 있는지, 어떤 문서와 정책이 현재 업무 기준으로 사용되는지를 구분해야 합니다. 오래된 문서, 중복 자료, 출처가 불분명한 정보가 함께 연결되면 AI 답변의 품질은 오히려 낮아질 수 있습니다.

기술적으로는 검색증강생성(RAG), 벡터 검색, 지식 그래프, 권한 관리, 출처 추적 등이 맥락 구축의 주요 기반이 됩니다. 그러나 핵심은 기술 자체가 아니라, 조직이 보유한 지식의 선별·관리 기준에 있습니다. AI가 참고할 수 있는 정보의 범위와 신뢰 수준이 명확해야 답변의 정확성과 설명 가능성도 높아집니다.

맥락 계층은 AI 시스템이 단순히 데이터를 검색하는 수준을 넘어, 조직 내 지식의 관계와 우선순위를 반영하도록 합니다. 표준 운영 절차, 정책 문서, 고객 이력, 과거 의사결정 기록, 검증된 전문 지식이 연결될 때 AI는 조직의 실제 판단 기준에 더 가까운 답변을 제공할 수 있습니다.

다만 맥락 계층이 복잡해질수록 오래된 문서, 중복 자료, 권한이 맞지 않는 정보가 함께 노출될 수 있습니다. 따라서 AI가 참조할 지식은 지속적으로 정리되어야 하며, 정보의 최신성, 출처, 권한 기준도 함께 관리되어야 합니다.

왜 지금 중요한가

AI 활용이 확대될수록 기업 고유의 맥락은 AI 성과를 좌우하는 핵심 요소가 됩니다. 기업이 AI에 제공할 수 있는 내부 지식의 품질과 관리 수준에 따라, AI 답변의 정확성과 신뢰도가 달라집니다.

맥락이 부족한 AI는 일반적인 답변은 할 수 있지만, 자사의 전략과 고객 약속, 리스크 기준에 맞는 판단을 내리기는 어렵습니다. 반대로 내부 지식이 체계적으로 연결된 AI는 조직의 업무 기준에 맞춰 더 정확하고 일관된 답변을 제공할 수 있습니다.

따라서 기업은 지식을 단순히 축적하는 데 그치지 않고, AI가 활용할 수 있는 형태로 정리하고 관리해야 합니다. 데이터, 문서, 정책, 의사결정 이력을 하나의 지식 체계로 연결하는 것이 AI 활용의 정확성과 신뢰성을 높이는 핵심 조건입니다.

2026년에 변화하는 것

2026년에는 AI가 참고할 모델 성능만큼이나, 기업 내부 지식을 얼마나 정확하고 최신 상태로 관리하느냐가 중요해지고 있습니다:

- **지식 체계 관리:** 지식은 더 이상 단순히 저장되는 자료가 아니라, AI 답변과 업무 판단을 뒷받침하는 핵심 자산으로 관리됩니다.
- **출처 기반 답변 강화:** AI가 생성한 결과에는 참고한 문서와 근거가 함께 제시되고, 신뢰할 수 있는 정보에 기반한 답변이 중요해집니다.
- **선별된 지식 계층 구축:** 중요도와 신뢰도에 따라 문서를 정리하고, AI가 참조할 수 있는 정보의 범위와 우선순위를 명확히 합니다.
- **최신성 관리 강화:** 데이터와 문서의 출처, 최신성, 신뢰도를 함께 관리해 오래된 정보나 잘못된 정보가 AI 답변에 반영되지 않도록 합니다.

파운데이션 모델은 공개된 지식에는 강하지만, 기업 내부의 문서와 업무 기준에는 기본적으로 접근할 수 없습니다.

방대한 일반 지식을 설명할 수 있어도, 기업의 전략과 의사결정 기준을 스스로 이해할 수는 없습니다. 결국 맥락은 범용 AI의 지능을 조직의 실제 판단과 연결하는 핵심 조건입니다.

한국의 상황

한국은 공공데이터 개방, AI 학습용 데이터 구축, 데이터 바우처 지원 등을 통해 AI 활용 기반을 넓혀 왔습니다. 그러나 기업의 AI 성과는 공공 데이터의 활용 자체보다 내부 지식과 업무 기준을 정확하게 연결하고 관리하는 역량에 달려 있습니다.

국내 기업은 RAG 도입과 내부 문서 검색, 업무 시스템 연계를 확대하고 있지만, 오래된 문서와 중복 자료, 관리되지 않은 콘텐츠가 함께 연결되는 경우도 적지 않습니다. 앞으로는 데이터의 양보다 품질과 최신성, 출처, 권한 기준을 관리해 AI 답변의 정확성과 신뢰성을 높이는 것이 중요합니다.

리더십 우선순위

맥락 계층을 구축하기 위해서는 기술적 기반과 운영 체계를 함께 정비해야 합니다:

- **신뢰할 수 있는 지식 기반 구축:** 주요 업무 흐름에 필요한 문서, 정책, 데이터, 의사결정 기준을 식별하고 정비해야 합니다.
- **검색 체계 정비:** 정형 데이터, 문서 저장소, 정책 자료를 하나의 검색 체계로 연결해 AI가 승인된 정보에 기반해 답변하도록 해야 합니다.
- **출처와 최신성 관리:** AI가 참조하는 정보의 출처, 최신성, 권한 기준을 명확히 하고, 오래된 정보가 답변에 반영되지 않도록 관리해야 합니다.

변화의 신호

- **맥락 계층의 확장:** 기존 지식 기반을 검색과 요약에 그치지 않고, 새로운 에이전트와 업무 흐름에 연결합니다.
- **공식 자료 중심의 운영:** AI 답변이 검증된 내부 자료와 승인된 정책을 기반으로 생성됩니다.
- **최신 정보 유지 체계 확립:** 사용자와 관리자가 AI 답변의 출처와 갱신 상태를 확인할 수 있고, 정보 갱신 기준이 운영 절차에 반영됩니다.
- **일관된 답변 품질:** 고객 대응, 내부 보고, 의사결정 지원에서 AI 답변의 정확성과 일관성이 높아집니다.

AI 리더의 사고방식

리더는 맥락을 단순한 데이터 관리 문제가 아니라, AI를 조직의 실제 업무 기준과 연결하는 핵심 운영 과제로 봐야 합니다. 중요한 것은 조직의 지식이 최신 상태로 유지되고, 올바른 권한과 기준에 따라 AI에 제공되도록 관리하는 체계입니다.

AI가 신뢰할 수 있는 결과를 내기 위해서는 모델 성능만으로 충분하지 않습니다. 어떤 정보를 근거로 삼는지, 그 정보가 최신인지, 누가 검증했는지, 어떤 기준에 따라 활용되는지가 함께 관리되어야 합니다. 맥락은 AI 답변을 조직의 판단 기준에 맞추고, 신뢰할 수 있는 의사결정으로 연결하는 기반입니다.

시뮬레이션 환경

시뮬레이션 환경은 현실의 업무 환경을 디지털로 재현해 AI를 안전하게 학습·검증하는 환경을 의미합니다. 디지털 트윈, 합성 데이터, 시뮬레이션 랩 등을 활용해 실제 환경에서 검증하기 어렵거나 위험한 상황을 사전에 점검할 수 있습니다.

AI 활용 범위는 점점 리스크가 큰 핵심 업무 영역으로 확대되고 있습니다. 금융 거래 처리, 의료 임상, 실물 자산 제어 등 오류가 곧 치명적인 손실로 이어지는 영역에서는 운영 환경에서의 시행착오 자체가 막대한 비용과 리스크를 동반합니다.

시뮬레이션 환경은 이러한 위험을 줄이기 위한 대안입니다. 실제 업무와 운영 조건을 디지털 공간에 재현해, 고객과 직원, 자산이 영향을 받기 전에 AI 업무 흐름을 시험하고 예외 상황과 통제 장치를 검증할 수 있도록 합니다. 디지털 트윈, 시뮬레이션 랩, 합성 데이터 파이프라인은 안전이 중요한 산업에서 활용해 온 검증 방식을 AI 영역으로 확장합니다.

이는 AI 시스템의 자율성이 높아질수록 더욱 중요해집니다. 에이전트 기반 업무 흐름은 단순히 한 번 답변하고 끝나는 것이 아니라, 일정 시간 동안 계획을 세우고 행동하며 여러 시스템과 상호작용하기 때문에 그만큼 예상치 못한 동작이 발생할 가능성도 커집니다. 시뮬레이션 환경에서는 실패 상황은 물론, 시스템 간의 업무 인계, 복구 절차를 포함한 전체 흐름을 사전에 재현하고 점검할 수 있습니다.

기업이 가질 수 있는 전략적 이점은 속도를 높이면서도 운영 리스크를 낮출 수 있다는 데 있습니다.

시뮬레이션은 팀이 수천 개의 시나리오를 빠르게 검토하고, 극단적 상황을 안전하게 검증하며, 드물게 발생하는 사건이 실제로 발생하기를 기다리지 않고 정책과 임계값을 조정할 수 있습니다. 또한 실제 데이터가 부족하거나 민감한 정보를 포함한 경우, 또는 데이터 주권 및 개인정보 보호 규제로 활용이 제한되는 경우에는 시뮬레이션 데이터가 유용한 대안이 됩니다. 결과적으로 조직은 AI 시스템의 안정성을 높이고, 출시 전 검증 근거를 확보하는 실용적인 방법이 될 수 있습니다.

왜 지금 중요한가

시뮬레이션 환경은 AI 시스템의 신뢰성과 도입 속도를 동시에 높이는 기반이 됩니다. 업무 흐름이 자동화되고 에이전트가 여러 도구와 시스템을 넘나들게 되면서, 기존 테스트 방식만으로는 실제 시나리오의 복잡성을 충분히 검증하기 어렵습니다. 특히 데이터가 제한적이거나 중대한 사건이 드물게 발생하는 영역에서는 사전에 다양한 조건을 재현하고 통제 효과를 확인할 수 있는 환경이 필요합니다.

2026년에 변화하는 것

AI 네이티브 기업은 단순히 모델을 점검하는 수준을 넘어, 실제 운영 조건에 가까운 환경에서 전체 시스템을 검증하는 방식으로 전환하게 됩니다.

- **디지털 리허설:** 팀이 전체 업무 흐름, 시스템 간 의존성 및 실패 가능성을 사전에 점검합니다.
- **합성 데이터 활용:** 실제 데이터가 부족하거나 활용이 제한될 때, 안전하고 대표성 있는 데이터셋으로 테스트를 보완합니다.
- **지속적인 검증 근거 확보:** 출시 과정에서 필요한 검증 자료와 판단 근거가 자동으로 축적됩니다.
- **강화된 통제 체계:** 실제 운영에 적용하기 전 스트레스 상황에서도 에스컬레이션, 예외 처리 및 복구절차가 제대로 작동하는지 검증합니다.

시뮬레이션 환경은 ‘추측과 기대에 의존’하던 검증 방식을 ‘실제 검증과 확신’에 기반한 접근으로 전환시킵니다.

한국의 상황

안전성과 데이터 보안이 중요한 국내 환경에서 시뮬레이션은 AI 검증을 위한 현실적인 대안으로 주목받고 있습니다. 정부는 국토교통부의 디지털 트윈 관련 사업과 자율주행 가상 시험 환경을 고도화하며, 실제 환경에 위험을 발생시키지 않고 다양한 시나리오를 검증할 수 있는 기반을 확대하고 있습니다.

산업 현장에서도 디지털 트윈을 통해 설계와 운영을 사전에 검증하고 있습니다. 조선·반도체 분야에서는 가상 공정 시뮬레이션을 통해 운영 조건을 미리 점검하고, 로봇 분야에서는 데이터 부족과 개인정보 제약을 보완하기 위해 합성 데이터를 활용하고 있습니다. 결국 시뮬레이션 환경은 AI가 실제 환경에서 안정적으로 작동할 수 있는지 확인하고, 기업이 속도와 안전성을 함께 확보하도록 돕는 핵심 역량이 되고 있습니다.

리더십 우선순위

- **시뮬레이션 랩 구축:** 우선순위가 높은 핵심 업무 흐름부터 시뮬레이션으로 검증할 수 있는 환경을 마련해야 합니다.
- **합성 데이터 활용 체계 마련:** 실제 데이터가 부족하거나 활용이 제한되는 경우, 현실성 있는 테스트 데이터를 생성·관리할 수 있는 체계를 구축해야 합니다.
- **근거 기반 출시:** 시나리오 검증 범위와 통제 장치의 성능을 출시 판단의 주요 기준으로 삼아야 합니다.

변화의 신호

- **충분한 사전 검증:** AI 시스템은 드물게 발생하는 예외 상황과 실패 가능성까지 검토한 뒤 운영 환경에 투입됩니다.
- **명확한 검증 근거:** 규제 기관과 리스크 관리 책임자가 별도 조사나 일회성 보고 없이도 통제 체계가 효과적으로 작동하고 있다는 근거를 확인할 수 있습니다.
- **운영 리스크 감소:** 시뮬레이션을 통해 다양한 조건을 반복적으로 점검함으로써 실제 운영 중 발생할 수 있는 사고의 빈도와 영향을 줄일 수 있습니다.
- **민감 데이터 사용 최소화:** 합성 데이터를 활용하면 실제 고객 데이터에 반복적으로 접근하지 않고도 테스트와 개선을 수행할 수 있습니다.

AI 리더의 사고방식

리더는 시뮬레이션 환경을 안전성과 속도를 동시에 확보하기 위한 핵심 역량으로 인식해야 합니다.

운영에 대한 확신은 시스템이 스트레스 상황과 예외 조건에서도 안정적으로 동작한다는 객관적인 검증 결과에서 나옵니다. 따라서 AI 리더는 중요한 업무 흐름에 충분한 사전 검증 체계가 마련되어 있는지, 통제 장치가 시뮬레이션 환경에서 안정적으로 작동하는지, 출시 과정에서 필요한 검증 근거가 지속적으로 축적되는지에 집중해야 합니다.

디지털 트윈(Digital Twins): 물리적 자산, 프로세스, 또는 시스템을 정밀하게 재현한 가상 복제 모델로, 다양한 상황을 시뮬레이션하고 결과를 예측하는 데 활용된다.

AI 운영(AI Ops)

AI 운영(AI Ops)은 모델, 프롬프트, 에이전트의 전 과정을 체계화하여, 일상 업무 내 AI의 품질·비용·통제를 일관되게 관리하는 핵심 운영체계입니다.

AI 활용이 업무 전반으로 확대되면 AI는 더 이상 개별 프로젝트의 산출물에 머물지 않습니다. 조직 운영의 일부로 자리 잡게 되며, 그에 따라 리더가 관리해야 할 기준도 달라집니다. 이제 관건은 AI 시스템이 핵심 업무 시스템과 같은 수준의 신뢰성, 투명성, 복구 능력을 갖추고 안정적으로 운영되는지를 확인하는 것입니다.

AI 도입 초기에는 많은 조직이 별도의 운영 체계 없이 소규모 팀의 수작업에 의존했습니다. 프롬프트 조정, 모델 배포, 대시보드 점검, 오류 대응까지 대부분 개별 담당자가 처리하는 방식이었습니다. 그러나 활용 범위가 넓어질수록 이러한 방식은 확장에 한계를 보입니다. 데이터, 사용자 행동, 정책 설정, 모델 버전의 작은 변화만으로도 결과가 달라질 수 있기 때문입니다. 비용은 통제되지 않은 채 증가할 수 있으며, 장애로 드러나지 않더라도 품질은 점진적으로 저하될 수 있습니다.

AI 운영(AI Ops)은 AI를 안정적으로 운영하기 위한 관리 체계입니다. 여기에는 모델과 프롬프트를 버전별로 관리하는 레지스트리, 업데이트를 사전에 검증하고 안정적으로 배포하기 위한 CI/CD(지속적 통합·배포), 운영 과정에서 나타나는 결과 편차와 품질 저하, 반복적인 실패 패턴을 점검하는 모니터링 체계가 포함됩니다. 또한 AI 사용량이 늘어날수록 연산 비용도 빠르게 증가할 수 있으므로, 비용을 관리하고 통제하기 위한 기준과 장치가 필요합니다.

이는 단순히 새로운 도구를 도입하는 데서 그치지 않습니다. AI 관련 문제는 일반적인 운영 장애와 같은 수준에서 관리되어야 합니다. 문제가 발생했을 때 누가 무엇을 판단하고 조치할지, 어떤 절차에 따라 대응할지, 필요할 경우 어떻게 이전 상태로 되돌릴지를 미리 정해 두어야 합니다. 또한 기술적 운영 데이터를 사업 지표와 연결하는 것도 중요합니다. 경영진은 고객 성과, 규정 준수, 재작업 등 자신에게 중요한 측면에서 모델이 제대로 작동하고 있는지 파악할 수 있어야 합니다.

AI 운영(AI Ops)은 조직 내 역할과 책임을 명확히 하는 데도 필요합니다. 모델 배포, 모니터링, 권한 관리, 비용 관리 등 공통 운영 기반은 IT·데이터 전담 조직이 관리하고, 현업 부서는 이를 활용해 업무 성과를 높이는 데 집중해야 합니다.

운영 관리와 현업 활용의 역할이 구분되면 안정성을 유지하면서도 AI 적용 속도를 높일 수 있습니다. 이를 통해 AI 운영은 알고리즘의 가능성을 실제 사업 성과로 연결하는 기반이 됩니다.

왜 지금 중요한가

AI 에이전트가 여러 업무 시스템과 연계되어 작업을 수행하기 시작하면서, 작은 오류도 고객 경험 저하나 규정 준수 리스크로 이어질 가능성이 커졌습니다. 이제 AI의 실패는 시스템이 완전히 멈추는 형태가 아니라, 사업 성과가 흔들리는 방식으로 나타날 수 있습니다. 따라서 리더는 변경 사항을 통제하고, 성능을 지속적으로 점검하며, 문제가 발생했을 때 빠르게 복구할 수 있는 운영 체계를 갖춰야 합니다. AI 운영(AI Ops)은 이를 가능하게 하는 핵심 관리 체계입니다.

2026년에 변화하는 것

AI는 개별 모델이나 프로젝트 산출물이 아니라, 지속적으로 운영하고 개선해야 하는 업무 서비스로 관리됩니다. 이에 따라 배포, 모니터링, 사고 대응, 복구 절차에 대한 기준이 더 중요해집니다.

- **레지스트리와 배포 기준:** 모델, 프롬프트, 정책을 버전별로 관리하고, 테스트와 승인 절차를 거쳐 운영 환경에 적용합니다.
- **비즈니스 관점의 모니터링:** 기술적 상태뿐 아니라 고객 성과, 운영 성과, 품질 변화, 실패 양상까지 함께 점검합니다.
- **사고 대응 체계 정비:** AI 오류와 장애는 담당 역할, 대응 절차, 롤백 기준, 사후 개선 체계를 갖춘 운영 이슈로 관리합니다.

정교한 AI 모델이라도 조직 전반에서 일관되게 운영되지 못하면 실질적인 사업 가치로 이어지기 어렵습니다.

한국의 상황

한국은 생성형 AI 도입은 빠르게 진행되고 있지만, 이를 안정적으로 운영·관리하는 체계는 아직 미흡합니다. 상당수 기업은 모델 배포와 관리를 개별 담당자나 특정 팀의 수작업에 의존하고 있어, 전사 차원에서 일관된 성능과 품질을 유지하는 데 한계를 보이고 있습니다.

이러한 격차는 AI를 실제 업무에 적용하고 운영할 전문 인력과 관리 기반이 부족한 데서 비롯됩니다. 최근에는 배포와 모니터링을 자동화하고, 품질 저하나 반복적인 오류를 조기에 감지하기 위한 운영 체계 도입이 늘고 있습니다. 그러나 문제 발생 시 책임 소재를 구분하고 신속하게 대응하는 체계는 여전히 취약합니다. 국내 기업은 시연 중심의 AI 활용을 넘어, 안정성과 신뢰성을 확보할 수 있는 운영 역량을 갖춰야 합니다.

리더십 우선순위

- **운영 기준 수립:** 품질, 결과 편차, 응답 지연 시간, 성과 단위당 비용에 대한 관리 기준을 정하고, AI가 적용되는 주요 업무가 배포 전후에 이 기준을 충족하도록 해야 합니다.
- **명확한 사업 책임 부여:** 승인, 모니터링, 사고 대응, 원상복구에 대한 역할과 책임을 구분하고, 운영 중인 모든 모델과 에이전트에 대해 성과와 리스크를 책임질 최종 책임자를 지정해야 합니다.
- **AI 사고 대응 체계 마련:** 주요 활용 사례별로 대응 절차, 보고·전파 경로, 원상 복구 방안을 마련하고 정기적으로 점검해야 합니다. 이를 통해 AI의 응답이나 처리 결과가 예상과 달라지는 경우에도 조직이 신속하게 대응할 수 있어야 합니다.

변화의 신호

- **변경 관리 체계:** 모델과 프롬프트, 정책 업데이트가 표준화된 절차에 따라 배포되며, 문제가 발생하면 신속하게 이전 상태로 복구할 수 있습니다.
- **성능 가시성 확보:** 품질 저하와 드리프트 징후가 조기에 포착되고, 리더는 성능 변화와 대응 현황을 명확히 확인할 수 있습니다.
- **명확한 책임 소재:** 운영 중인 모델별 책임자와 장애 대응 역할이 지정되어, 문제가 발생했을 때 신속한 판단과 조치가 이루어집니다.
- **개선 체계 정착:** 장애 원인과 수정 사항이 기록·공유되어, 동일한 유형의 문제가 반복되지 않도록 운영 방식이 지속적으로 개선됩니다.

AI 리더의 사고방식

AI 운영(AI Ops)은 '기능 배포' 중심의 접근에서 '신뢰성 확보' 중심의 접근으로 전환하는 것입니다. 이 관점에서는 보여주기식 성과보다 안정적인 운영이 더 중요합니다. 문제를 조기에 발견하고, 성과를 정확하게 측정하며, 시스템을 지속적으로 개선하는 조직일수록 AI를 신뢰할 수 있는 자산으로 발전시킬 수 있습니다.

이 보고서에서는 AI 모델뿐 아니라 프롬프트, 에이전트, 정책, 비용, 사고 대응까지 포함하는 운영 체계를 'AI 운영(AI Ops)'으로 정의합니다.

AI 운영(AI Ops)은 모델의 학습·검증·배포·모니터링을 중심으로 하는 MLOps보다 넓은 개념이지만, 아직 그 범위와 정의는 정립되는 과정에 있습니다.

2023년 본격화된 컴퓨팅 자원 부족은 AI 산업에 하드웨어 제약의 현실을 분명히 인식시키는 계기가 되었습니다. 이후 컴퓨팅 용량은 크게 늘었지만, AI 수요가 더 빠르게 증가하면서 컴퓨팅 자원은 여전히 AI 산업의 구조적 병목으로 남아 있습니다.

AI 모델 운영에 필요한 GPU, AI 가속기, 클라우드 인프라 등 컴퓨팅 자원은 기업의 핵심 제약 요인으로 부상했습니다. AI가 실험 단계를 넘어 일상 운영으로 확산되면서, 모델 접근성보다 컴퓨팅 자원의 안정적 확보와 효율적 배치·관리가 더 중요한 과제가 되고 있습니다.

그동안 많은 조직은 컴퓨팅 자원을 필요할 때마다 사용하는 클라우드 기반 자원으로 인식해 왔습니다. 그러나 AI 활용 규모가 커지면서 수요 급증, GPU 공급 불균형, 비용 부담 확대가 나타났고, 컴퓨팅 자원 관리는 소비 중심에서 계획 중심으로 전환되고 있습니다.

컴퓨팅 자원 관리는 워크로드 배치, 용량 확보, 운영 효율화의 기준을 정하는 일입니다. 여기에는 엣지, 온프레미스, 국내 클라우드 리전의 선택, 예약 용량과 약정, 업무 특성에 따른 인프라 등급 배치 기준이 포함됩니다.

핵심은 단위 경제성입니다. 경영진은 작업 한 건을 처리하는 비용, 업무 사례별 토큰 수, 모델 재학습에 필요한 GPU 시간, 품질 저하에 따른 재작업 부담을 함께 확인해야 합니다. 이를 바탕으로 소형 모델로 충분한 업무에는 비용 효율적인 자원을 우선 적용하고, 필요성이 입증될 때 더 큰 자원을 투입해야 합니다.

컴퓨팅은 비용뿐 아니라 환경 측면에서도 관리가 필요합니다. AI 인프라의 에너지 사용량, 냉각 효율, 물 사용량이 중요한 지표로 떠오르면서 컴퓨팅 자원 관리는 인프라 구축을 넘어 사업성, 재무 영향, 운영 효율을 함께 고려하는 경영 의제가 되고 있습니다.

왜 지금 중요한가

많은 조직은 AI 기능을 성공적으로 검증한 뒤에도, 이를 대규모로 활용하는 단계에서 상당한 비용과 용량 제약에 직면하고 있습니다. 이에 따라 컴퓨팅 관련 결정은 제품 선택, 서비스 수준, 출시 속도에 직접적인 영향을 미치기 시작했습니다.

이러한 변화를 가장 먼저 체감하는 곳은 AI를 이용 빈도가 높은 고객 접점과 핵심 운영에 도입한 조직입니다. 이들 조직에서는 고객 문의 처리, 청구 심사, 사용자와의 상호작용이 늘어날 때마다 직접적인 컴퓨팅 비용과 용량 부담이 함께 증가합니다.

2026년에 변화하는 것

2026년의 변화는 컴퓨팅 자원의 확보와 활용 전반에 집중됩니다.

- **계약 기반 용량 확보:** 주요 공급업체와 GPU 플랫폼이 예약, 우선순위 계층, 장기 약정 중심으로 재편되면서 용량 계획은 기술적 고려를 넘어 사업 협상의 영역으로 옮겨갑니다.
- **공급업체 다변화:** 전문 GPU 공급업체, 지역별 옵션, 관리형 추론 서비스가 성숙하면서 단일 공급업체 의존에서 벗어나 컴퓨팅 포트폴리오를 다변화하는 전략이 요구됩니다.
- **워크로드 배치 정교화:** 지연 시간, 데이터 주권, 비용을 고려해 엣지, 온프레미스, 리전 전반에 워크로드를 의도적으로 나누어 배치하는 체계가 자리 잡습니다.
- **효율성의 전략적 부상:** 제한된 컴퓨팅 자원을 최대한 활용하기 위해 소형 모델 활용과 스마트 라우팅은 단순한 최적화를 넘어 필수 전략으로 자리 잡습니다.

컴퓨팅은 AI의 가능성을 현실의 비용으로 전환합니다.

한국의 상황

국내 기업에 컴퓨팅은 필요할 때 사용하는 단순 유틸리티가 아니라, 단위 경제성을 기준으로 관리해야 할 전략 자원입니다. 가격 변동성과 인프라 비용 상승에 대응해 국내 시장에서도 AI 업무의 특성과 사용 기간에 맞춰 용량을 계획적으로 확보하고 배치하는 방식이 자리 잡아가고 있습니다.

이러한 비용 압박은 국산 AI 반도체의 상용화를 앞당기는 요인으로 작용하고 있습니다. 고효율 국산 반도체가 공공 클라우드 등에 채택되어 검증을 거치고 있으며, '학습은 범용 가속기로, 추론은 국산 반도체로' 나누는 포트폴리오 전략도 주목받고 있습니다. 이제 컴퓨팅은 기술적 조정의 문제가 아니라, 재무 관점의 비용 관리가 필요한 전략 의제입니다.

리더십 우선순위

- **컴퓨팅 비용 관리 체계화:** AI 과제의 사업성 검토 단계에서 필요한 GPU·클라우드 용량과 예상 사용 비용을 함께 점검해야 합니다.
- **배치 기준 수립:** AI 업무의 중요도와 사용 패턴에 따라 사전 확보 자원과 수요에 따라 추가 확보할 자원을 구분해 적용해야 합니다.
- **이동성을 고려한 설계:** 인프라나 공급업체에 과도하게 종속되지 않도록 설계하고, 성능 요건을 충족하는 경우 소형 모델을 우선 활용해야 합니다.

컴퓨팅 비용과 환경 영향에 대한 관리 요구가 높아지면서, 리더는 AI 업무에 필요한 자원을 어디에, 얼마나 배치할지 더 신중하게 판단해야 합니다.

변화의 신호

컴퓨팅은 과거 비가시적 기술 비용으로 인식되던 데서 벗어나, 경영진이 직접 관리하고 설명할 수 있는 사업 자원으로 전환되고 있습니다.

- **성과 단위당 비용 가시화:** 컴퓨팅 지출이 작업, 처리 건수, 서비스 성과 단위로 관리됩니다.
- **핵심 업무의 자원 확보:** 우선순위가 높은 AI 업무에 필요한 GPU·클라우드 자원이 사전에 확보됩니다.
- **가용 용량 확대:** 모델 경량화와 자원 배치 개선을 통해 추가로 활용할 수 있는 용량이 늘어납니다.
- **환경 영향 관리:** 에너지 사용량과 탄소 배출 등 AI 인프라의 환경 영향이 보고·관리됩니다.

AI 리더의 사고방식

'성공적인' 파일럿도 사용량이 늘어나면 감당하기 어려운 서비스가 될 수 있습니다. 기술 리더는 어떤 AI 업무에 사전 확보한 자원을 배정해야 하는지, 어떤 업무는 수요에 따라 자원을 추가 확보해도 되는지, 어디에서 자원 활용도를 높여 여유 용량을 확보할 수 있는지 명확히 파악해야 합니다.

이제 컴퓨팅은 단순한 기술 조정의 문제가 아니라 전략적 예산 논의의 대상입니다. AI를 확장하려는 조직은 컴퓨팅 비용을 성과 단위로 관리하고, 자원 사용량을 지속적으로 점검해야 합니다.

설계 단계의 신뢰

신뢰는 AI가 조직의 의사결정과 업무 운영에 안정적으로 활용되기 위한 기본 조건입니다.

설명 가능성, 편향·공정성 관리, 데이터 스튜어드십은 AI 결과가 이해되고 검증되며 책임 있게 활용되도록 만드는 기준입니다. 여기에 지속적인 근거 관리, 모니터링, 통제 절차가 더해질 때 AI 신뢰는 일회성 검토에 그치지 않고 운영 체계 안에 자리 잡을 수 있습니다.

인간 중심의 통제는 AI가 실질적인 영향을 미치는 영역에서 사람의 검토와 이의제기, 개입 절차를 보장해 책임성과 수용성을 높이는 장치입니다.



책임 있는 AI

책임 있는 AI는 AI 활용 과정에서 발생할 수 있는 위험을 관리하면서, AI가 창출할 수 있는 가치를 안정적으로 실현하기 위한 운영 원칙과 실행 체계입니다. 단순히 위험을 줄이는 데 그치지 않고, 신뢰를 바탕으로 AI 활용 범위를 넓히기 위한 접근입니다.

AI 활용이 확대될수록 신뢰는 선택 사항이 아니라 기본 조건이 됩니다. 설명 가능성, 개인정보 보호, 편향성 관리, 공정성 확보, 유해성 관리, 사람의 관리·감독 체계는 AI를 지속적으로 활용하기 위한 필수 요건입니다. 이러한 요건은 조직이 AI의 결과를 이해하고 설명하며, 문제가 발생했을 때 책임 있게 대응하는 기반이 됩니다.

책임 있는 AI는 선언이나 원칙만으로 완성되지 않습니다. 명확한 역할과 책임, 위험 수준에 따른 승인 절차, 데이터 관리 기준, 모니터링과 사후 점검 체계가 실제 업무 절차에 포함되어야 합니다. 이런 운영 기반은 AI가 빠르게 확산되는 상황에서도 필요한 통제 수준을 유지하도록 합니다.

AI에 대한 신뢰가 흔들리면 도입 속도는 급격히 늦어집니다. 한 번의 오류나 개인정보 침해, 차별적 결과, 설명 불가능한 의사결정은 특정 시스템의 문제를 넘어 조직 전체의 AI 활용에 대한 불신으로 이어질 수 있습니다. 반대로 책임 있는 AI 체계가 마련되면 리더는 더 빠르게 의사결정을 내리고, 현장에 권한을 위임하며, 통제력을 유지한 상태에서 AI 활용 범위를 넓힐 수 있습니다.

PwC는 책임 있는 AI를 조직이 위험과 가치를 일관되고 투명하며 책임 있게 관리해 AI의 잠재력을 극대화하는 실행 체계로 봅니다. 이는 책임 있는 AI가 별도의 정책 문서가 아니라, 모델 개발과 업무 적용, 배포, 운영, 개선 전 과정에 적용되는 관리 방식이기 때문입니다. 모델과 업무 흐름이 바뀌더라도 AI 활용이 사업 목표, 고객 기대, 규제 기준에서 벗어나지 않도록 지속적으로 점검하는 체계가 필요합니다.

국내에서도 AI 윤리, 신뢰성, 개인정보 보호, 설명 가능성, 책임성에 대한 요구가 구체화되고 있습니다. 특히 생성형 AI와 에이전트형 AI가 업무 현장에 도입되면서 민감 정보 입력, 데이터 활용, 결과 검증, 책임 소재에 대한 관리 필요성이 커지고 있습니다. 많은 조직이 거버넌스 체계보다 빠른 속도로 AI를 도입하고 있는 만큼, AI 활용 확대와 관리 체계 사이의 격차를 줄이는 것이 시급합니다.

왜 지금 중요한가

책임 있는 AI는 더 이상 윤리 원칙이나 내부 지침의 문제가 아닙니다. 고객 응대, 직원 업무, 규제 대응, 의사결정 지원 등 실제 업무에 AI를 적용하려는 조직이 반드시 갖춰야 할 운영 요건입니다.

2026년에 변화하는 것

책임 있는 AI는 원칙 중심의 논의에서 벗어나 전사적 운영 체계로 전환됩니다. 빠른 배포와 변화하는 규제 환경, 높아지는 고객 기대에 대응하려면 AI를 기획·개발·배포·운영하는 전 과정에서 책임 있는 관리 기준을 적용해야 합니다.

- **거버넌스의 실행 체계화:** 정책과 요건은 별도 문서가 아닌 개발 표준, 출시 전 점검, 승인 절차, 책임 구조에 직접 반영됩니다. 책임 있는 AI는 원칙을 정하는 단계에서 끝나는 것이 아니라, 실제 업무 프로세스에서 구현되어야 합니다.
- **위험 등급 기준의 표준화:** AI 활용 사례의 위험 수준과 영향도를 기준으로 통제 강도를 달리 적용합니다. 저위험 업무는 신속하게 추진하되, 고위험 업무는 더 엄격한 검토와 승인 절차를 거치도록 합니다. 이를 통해 속도와 통제 사이의 균형을 맞출 수 있습니다.
- **신뢰의 운영 기준화:** 신뢰는 기업이 대외적으로 내세우는 구호가 아니라, 업무 과정에서 반복적으로 확인되는 기준이 됩니다. 명확한 책임 소재, 위험 수준에 맞는 통제, 배포 이후의 모니터링과 개선 절차가 일상적인 운영 방식으로 자리 잡아야 합니다.

조직이 AI를 안심하고 확장하려면 책임 있는 AI 체계를 먼저 갖추어야 합니다.

AI 활용의 가장 빠른 길은 절차를 생략하는 것이 아니라, 설계로 신뢰를 표준화해 승인을 신속히 하고 책임을 명확히 하며 근거를 자동화하는 것입니다. 따라서 책임 있는 AI는 더 빨리 출시하고 더 넓게 확장하도록 돕는 신뢰 확보 전략입니다.

한국의 상황

국내 AI 거버넌스는 정부의 AI 윤리 원칙과 관련 가이드라인을 기반으로, 실무에서 적용 가능한 통제 기준을 마련하는 방향으로 발전해 왔습니다. 이러한 흐름 속에서 국내 기업들도 선언적 수준의 윤리 지침을 마련하는 단계를 넘어, AI 리스크 관리 체계를 실제 사내 운영 시스템에 반영하는 단계로 나아가고 있습니다.

이에 따라 기업들은 형식적인 문서 검토에 머무르지 않고, 실제 운영 환경에서 AI의 성능과 위험 요인을 검증하는 방향으로 관리 체계를 고도화하고 있습니다. 또한 기획·개발 부서와 분리된 독립적 위험관리 조직을 두고, AI 활용 과정에서 발생할 수 있는 리스크를 사전에 점검하는 체계를 마련하고 있습니다. 이러한 통제 구조는 리더가 전사적 AI 도입을 보다 확신 있게 승인할 수 있도록 뒷받침하는 실질적 근거가 됩니다.

리더십 우선순위

- **책임을 고려한 설계:** 모든 AI 이니셔티브는 명확한 목적, 리스크 등급, 책임 주체를 정하는 것에서 출발해야 합니다. 의사결정 권한, 중단 및 진행 기준, 필요 시 시스템을 멈출 수 있는 기준도 사전에 합의되어야 합니다. 신뢰는 마지막 단계에서 점검하는 항목이 아니라, 무엇을 만들고 어떻게 운영하며 누가 책임질지를 결정하는 설계의 기준입니다.
- **업무 흐름에 내재화된 통제:** 신뢰는 별도의 절차가 아니라 실제 업무와 시스템 안에서 작동해야 합니다. 데이터 최소화, 편향성과 공정성 테스트, 개인정보 보호 설계, 파이프라인에 내장된 가드레일과 모니터링, 중요한 의사결정 지점에서의 사람 개입이 여기에 포함됩니다. 이러한 통제는 리스크 관리를 자연 요인이 아니라 AI 활용의 속도와 안정성을 높이는 기반으로 전환합니다.

- **확장에 맞는 근거 체계:** 신뢰는 선언이 아니라 근거로 입증되어야 합니다. 각 사용 사례는 정책, 평가 결과, 모니터링 내역 등 디지털 근거를 함께 갖추고, 글로벌 기준과도 정합성을 확보해야 합니다. 규제 기관은 적용 범위, 도입 현황, 사고 발생 여부 등 핵심 지표를 확인할 수 있어야 하며, 필요할 경우 이를 바탕으로 AI 운영의 적정성을 검증할 수 있어야 합니다.

변화의 신호

- **거버넌스 체계 정착:** AI 활용 과정에 필요한 통제 수단이 기본적으로 마련되어 있습니다.
- **출시 절차와의 연계:** 책임 있는 AI 기준이 구축과 출시 절차에 반영되어 있습니다.
- **책임과 권한의 명확화:** AI 리스크와 운영에 대한 책임 주체가 명확히 지정되어 있습니다.
- **의사결정 근거 확보:** 주요 의사결정의 근거를 필요할 때 제시할 수 있습니다.
- **수명 주기 전반의 관리:** 기획과 구축, 운영, 폐기 단계까지 책임 있는 AI 기준이 적용됩니다.

AI 리더의 사고방식

책임 있는 AI는 AI 확장을 가능하게 하는 전제 조건으로 다뤄져야 합니다. 리더는 원칙이 실제 업무와 운영 과정에서 실행되고 있는지 확인해야 합니다. 중요한 것은 책임 있는 AI를 선언하는 것이 아니라, 조직이 이를 일상적인 의사결정과 운영 체계 안에서 지속적으로 입증할 수 있는지입니다.

설명 가능성

신뢰할 수 있는 AI는 결과만 제시하지 않습니다. 판단의 근거를 설명하고, 문제가 있을 때 검토와 이의제기가 가능해야 합니다.

설명 가능성과 이의제기 가능성은 AI가 실제 의사결정과 업무 결과에 영향을 미치는 지점에서 특히 중요합니다. AI의 결과를 사람이 이해할 수 있는 판단 근거로 제시하고, 시스템이 맥락을 놓쳤거나 잘못된 규칙을 적용했거나 부정확한 정보에 의존했다고 판단될 때 이를 검토할 수 있는 공정한 절차가 마련되어야 하기 때문입니다.

실무적으로 설명 가능성은 사람의 검토를 지원하기 위해 어떤 근거와 증거, 불확실성 정보를 사용자와 의사결정권자, 감독 기관에 제공할 것인지 정하는 설계 과제입니다.

이러한 기능이 제품 경험과 출시 절차에 내재화되면, 조직은 영향력이 큰 업무 흐름에도 AI를 보다 안정적으로 적용할 수 있습니다. 반대로 이런 장치가 없으면 팀은 수작업 보완이나 비공식 보고 절차에 의존하게 됩니다. 그 결과 AI의 판단을 설명하거나 정당화하기 어려워지고, 시간이 지나도 체계적인 개선으로 이어지기 어렵습니다.

설명 가능성이 어려운 근본적인 이유는 현대 AI 아키텍처의 복잡성과 변화 가능성에 있습니다. 하나의 의사결정은 모델 응답, 데이터베이스 쿼리, 프롬프트 구성, 외부 도구 연동, 거버넌스 설정 등 여러 요소가 결합되어 만들어집니다. 또한 이러한 요소는 같은 사용자 요청에서도 시점과 조건에 따라 달라질 수 있습니다. 그러나 많은 기업은 특정 결과가 생성된 시점에 어떤 구성 요소와 설정, 데이터 출처가 함께 작동했는지를 정확히 추적할 수 있는 감사 체계를 충분히 갖추지 못하고 있습니다.

한편 AI 의사결정의 투명성을 요구하는 규제 흐름은 최종 결과 뿐 아니라, 그 결과를 만들어낸 개발·구현 과정까지 문서화할 것을 강조하고 있습니다. 이는 기업에 적지 않은 기술적·운영적 부담으로 작용합니다.

왜 지금 중요한가

책임 있는 AI를 구현하려면 설명 기능과 검토 절차가 중요한 의사결정 업무 흐름 안에 직접 포함되어야 합니다. 이를 통해 감독은 부담스러운 사후 조사에서 벗어나, 근거를 바탕으로 신속하게 검증하는 절차로 전환될 수 있습니다.

2026년에 변화하는 것

선도 기업들은 설명 가능하고 이의제기가 가능한 AI를 구축·운영하는 방식에서 다음과 같은 변화를 보이고 있습니다.

- **설명 기능의 내재화:** 영향력이 큰 의사결정에는 점차 간결한 판단 근거, 주요 입력값, 불확실성에 대한 명확한 설명이 함께 제공되고 있습니다. 이를 통해 사용자는 결과를 그대로 수용할지, 추가 검토나 상위 절차로 넘길지 판단할 수 있습니다.
- **기록 체계의 고도화:** 버전이 관리되는 모델과 프롬프트, 검색된 자료, 도구 실행 내역이 업무 흐름의 일부로 기록되고 있습니다. 이를 통해 팀은 특정 결과가 어떻게 생성되었는지를 보다 빠르고 일관되게 재구성할 수 있습니다.
- **이의제기 절차의 공식화:** 검토와 이의제기는 영향력이 큰 의사결정에 필요한 표준 절차로 자리 잡고 있습니다. 이 과정에는 명확한 책임자, 처리 기한, 수정 경로가 함께 포함됩니다.

설명 가능성은 AI 제품 설계의 일부로 자리 잡고 있습니다. 민감한 서비스에서는 결과의 근거와 변화 내용, 이의제기 방법을 명확히 안내해야 합니다. 다만 이러한 설명과 검토 장치가 실제 위험 수준에 맞게 설계되어 있는지 점검하는 일은 여전히 리더의 과제로 남아 있습니다.

설명할 수 없는 의사결정은 확장할 수 없는 의사결정입니다.

한국의 상황

설명 가능성과 이의제기 가능성은 금융·의료처럼 AI 판단이 개인의 권리, 기회, 비용, 건강에 직접적인 영향을 미치는 영역에서 필수적입니다. 금융 서비스에서는 AI가 자격 심사, 가격 책정, 보험금 청구 분류, 부정행위 탐지 등에 활용되며 고객의 거래 조건과 서비스 접근성에 영향을 미치고 있습니다. 의료 영역에서도 AI 기반 의사결정 지원 도구가 임상 판단과 책임 체계 안으로 들어오고 있습니다.

AI 판단이 실제 삶의 조건에 영향을 줄수록, 고객과 환자, 규제 당국은 결과의 근거와 검토 가능성을 요구하게 됩니다. 판단의 이유를 알 수 없거나 오류를 바로잡을 절차가 없다면 공정성 논란과 책임 공백이 발생할 수 있기 때문입니다. 따라서 설명 가능성과 이의제기 가능성은 AI 의사결정의 투명성과 신뢰를 유지하기 위한 핵심 장치입니다.

리더십 우선순위

대부분의 조직에서 시작점은 어떤 의사결정이 가장 중요한지 파악하는 것입니다.

- **주요 의사결정 지점 식별:** 설명이 필요한 고객, 직원, 규제 관련 의사결정 지점을 먼저 파악해야 합니다. 이후 의사결정 유형별로 어떤 수준의 설명이 필요한지 명확한 기준을 세워야 합니다.
- **검토 절차의 체계화:** 영향력이 큰 업무 흐름에는 이의제기와 재검토 절차가 포함되어야 합니다. 이를 위해 책임자, 필요한 근거, 처리 기한, 수정 경로를 사전에 명확히 정의해야 합니다.
- **의사결정 기록 관리:** 시스템이 결과에 영향을 미친 입력값, 모델과 프롬프트 버전, 참조한 자료를 기록하도록 해야 합니다. 이를 통해 조직은 중요한 의사결정이 어떻게 이루어졌는지를 일관되게 재구성하고 설명할 수 있습니다.

변화의 신호

- **이해 가능한 근거 제공:** 사용자는 의사결정 시점에 판단의 논리와 핵심 근거를 확인할 수 있습니다.
- **결과 재구성 가능:** 중요한 의사결정이 어떤 과정을 거쳐 이루어졌는지 기록을 통해 신속하게 확인할 수 있습니다.
- **실질적인 이의제기 절차:** 검토와 재판단 절차가 정해진 기준과 기한 안에서 실제로 운영됩니다.

AI 리더의 사고방식

리더는 AI의 판단이 고객, 직원, 환자 등 이해관계자에게 어떤 영향을 미치는지 확인하고, 그 결과와 근거가 당사자에게 이해 가능한 방식으로 설명될 수 있어야 합니다. 또한 조직은 해당 판단이 어떤 입력값과 기준, 자료를 바탕으로 이루어졌는지 재구성할 수 있어야 합니다. 이러한 기반이 갖춰져야 AI 시스템이 확장되더라도 책임성과 서비스 품질, 신뢰를 유지할 수 있습니다.

편향성, 공정성, 유해성

18

편향성, 공정성 및 유해성 관리는 AI가 사람들의 삶을 향상시킬지, 아니면 소외시킬지를 결정합니다. 리더는 불평등한 결과를 예방하고 바로잡아야 할 책임이 있습니다.

AI의 편향과 불공정은 대개 명백한 악의로 드러나지 않습니다. 오히려 당사자가 납득하기 어려운 결과의 반복적인 패턴으로 나타납니다. 충분한 신용 이력이 있음에도 고위험군으로 분류되는 고객, 유사한 역량을 갖추고도 탈락하는 지원자, 실제 필요와 다른 대리 지표를 기준으로 우선순위가 밀리는 환자, 과거 데이터의 편향 때문에 불리한 평가를 받는 이용자가 그 예입니다.

AI로 인한 유해는 불평등한 대우에만 그치지 않습니다. 고정 관념을 강화하거나 비하적 결과물을 만들어 개인의 존엄성을 훼손할 수 있고, 합리적으로 이용할 수 있어야 할 서비스에 대한 접근을 제한해 기회를 박탈할 수도 있습니다. 또한 자신의 삶에 영향을 미치는 의사결정을 신뢰하지 못하게 만들어 제도 와 서비스 전반에 대한 신뢰를 약화시킬 수 있습니다.

AI가 업무 흐름에 깊이 편입될수록 이러한 부정적 영향은 빠르게 확산될 수 있습니다. 전체 성과 지표가 좋아 보이더라도 특정 집단이나 상황에서는 불공정한 결과가 반복될 수 있기 때문입니다. 따라서 편향과 공정성 관리는 사후 대응이 아니라, 설계와 운영 전반에 걸쳐 체계적이고 예방적으로 이루어져야 합니다.

기술적으로 중요한 것은 특정 의사결정에서 무엇을 '공정하다'고 볼 것인지 명확히 정의하고, 집단 간 불균형한 영향을 점검하며, 시스템이 기존의 불이익을 확대하지 않도록 업무 흐름을 설계하는 것입니다. 가장 효과적인 접근은 많은 조직이 예상하는 것보다 이른 단계에서 시작됩니다. 운영 전에 위험을 예방하는 것이 불만 제기, 재작업, 출시 후 조사로 문제를 발견하는 방식보다 더 안정적입니다.

왜 지금 중요한가

많은 조직은 AI 시스템에서 편향이 얼마나 쉽게 발생하는지 과소평가합니다. 편향은 명백한 오류보다 발견하기 어렵고, 학습 데이터, 설계 선택, 운영 과정 전반에 숨어 있을 수 있습니다. 과거 데이터의 격차나 편향된 지표, 접근하기 쉬운 정보만을 우선시하는 방식은 특정 집단에 불리한 결과로 이어질 수 있습니다. 이러한 결과가 반복되면 피드백 루프가 형성되어 시간이 지날수록 불평등이 강화될 위험도 있습니다.

공정성은 집단 간 평균 성과만으로 판단하기 어렵습니다. 의사 결정 유형, 업무 흐름, 예외 상황까지 함께 살펴야 실제 영향을 확인할 수 있습니다. 따라서 공정성 관리는 사후에 문제를 발견해 수정하는 방식에서 벗어나, 설계 단계부터 편향과 유해성을 예방하는 방향으로 전환되어야 합니다.

2026년에 변화하는 것

공정성은 점차 의사결정 시스템에 사후적으로 덧붙이는 기준이 아니라, 설계 단계부터 반영해야 하는 핵심 속성으로 다뤄지고 있습니다. 이에 따라 설계와 출시 단계에서 편향과 유해성을 예방하기 위한 장치가 함께 구축되고 있습니다.

- **출시 전 공정성 검증:** 중요한 의사결정 시스템에는 공정성 테스트와 위해 시나리오 검토가 출시 기준으로 포함됩니다. 팀은 집단별·의사결정별 격차와 오류율을 함께 점검합니다.
- **업무 절차에 반영된 통제:** 영향력이 큰 의사결정에는 임계값 정책, 사람의 검토 지점, 에스컬레이션 규칙 등 불공정한 결과를 줄이기 위한 통제가 포함됩니다.
- **지속적인 편향 개선:** 편향이 확인되면 데이터와 정책을 조정하고, 필요한 경우 업무 흐름을 재설계합니다. 공정성 관리는 일회성 점검이 아니라 지속적인 운영 과제입니다.

공정성은 AI 도입만으로 확보되지 않습니다.

평균 성과만으로는 특정 집단이나 상황에서 발생하는 격차를 확인하기 어렵습니다. AI 시스템의 편향과 유해를 줄이려면 공정성을 설계 단계부터 반영해야 합니다. 중요한 것은 편향이 발생할 가능성을 부정하는 것이 아니라, 운영 과정에서 편향이 확산되기 전에 이를 포착하고 바로잡을 수 있는 체계를 갖추는 것입니다.

한국의 상황

한국의 국가 AI 전략은 글로벌 기술 경쟁에 대응하기 위해 'AI 3대 강국(G3) 도약'을 핵심 목표로 제시하고 있습니다.

이 과정에는 고위험 AI의 위험을 모니터링하고, AI 안전 기술을 검증·공유하기 위한 AI 안전연구소 설립도 포함되어 있습니다.

실무적으로 이는 플랫폼 알고리즘의 공정성, AI 채용과 금융 대출 심사의 투명성, 공공 행정 서비스의 신뢰성처럼 국민의 일상과 직접 연결되는 의사결정 영역에서 편향을 최소화해야 한다는 의미를 갖습니다. 동시에 혁신을 위축시키지 않으면서도, AI가 공정하고 신뢰할 수 있는 방식으로 활용되도록 관리하는 것이 중요한 과제로 떠오르고 있습니다.

리더십 우선순위

- **의사결정별 공정성 기준 설정:** 중요한 의사결정에서 무엇을 불공정한 결과로 볼 것인지, 어떤 집단을 모니터링해야 하는지, 어느 수준의 격차를 허용할 수 있는지 사전에 합의해야 합니다.
- **출시 전 예방 조치 반영:** 영향력이 큰 의사결정 시스템은 출시 전 공정성·위해성 점검을 거쳐야 합니다. 이때 검토 결과와 판단 근거를 남겨 신속한 검증이 가능하도록 해야 합니다.
- **신속한 시정 체계 마련:** 편향 신호가 감지되면 데이터, 정책, 업무 절차를 적시에 조정할 수 있어야 합니다. 이를 위해 책임 주체와 대응 절차를 명확히 정해 두는 것이 중요합니다.

변화의 신호

- **사전 점검 정착:** 공정성 점검이 출시 과정의 일부로 운영됩니다.
- **집단별 결과 가시화:** 평균값이 아니라 집단별 결과를 기준으로 영향을 파악합니다.
- **신속한 시정:** 격차가 확인되면 책임 있게 원인을 파악하고 조정합니다.
- **운영 안정성 유지:** 시스템이 변경되거나 업데이트되더라도 공정성 기준이 유지됩니다.

AI 리더의 사고방식

공정성은 AI 확산 과정에서 리더가 의도적으로 관리해야 할 핵심 책임입니다. 리더는 AI 시스템이 사람들을 공정하고 존중하는 방식으로 대하고 있는지 지속적으로 확인해야 합니다.

이는 누가 혜택을 받고 누가 불이익을 받는지, 평균 성과 뒤에 어떤 격차나 위해가 숨어 있는지 적극적으로 살핀다는 의미입니다.

또한 조직은 불균형한 결과를 조기에 발견하고, 방어적으로 대응하기보다 신속하게 바로잡을 수 있는 문화를 갖춰야 합니다. AI 시스템이 변화하고 확장되더라도 공정성이 일상적인 의사결정과 운영 절차 안에서 지속적으로 관리될 때, 조직은 신뢰를 유지하면서 AI 활용을 확대할 수 있습니다.

데이터 스튜어드십

데이터 스튜어드십은 데이터를 책임 있게 관리해야 할 자산으로 보고, 수집부터 폐기까지 사용 기준과 책임을 정하는 체계입니다. 이를 통해 AI 수명 주기 전반에서 개인정보와 민감 정보가 안전하게 활용되도록 관리합니다.

AI는 데이터의 활용 가치를 크게 높입니다. 프롬프트, 검색 과정, 학습 데이터, 로그 파일 등은 정보를 빠르게 새로운 역량으로 전환할 수 있게 합니다. 그러나 동시에 민감 정보가 조직의 정책과 통제 체계보다 빠르게 의도치 않은 경로로 이동할 위험도 커집니다.

데이터 스튜어드십은 AI 활용이 법적 기준에 부합하고, 설명 가능하며, 개인과 조직의 권리를 존중하는 방식으로 이루어지도록 하는 기반입니다. 여기에는 개인정보 보호, 동의와 목적 제한, 데이터 최소화, 비식별화, 데이터 출처와 이동 경로의 추적 가능성 확보가 포함됩니다. 또한 AI 시스템이 초안 작성, 요약, 검색 과정에서 독점 정보나 영업 기밀을 의도치 않게 노출하지 않도록 보호하는 전략도 필요합니다.

데이터 관리가 어려운 이유는 민감 데이터가 활용 맥락에 따라 서로 다른 위험을 갖기 때문입니다. 예를 들어 서비스 제공에는 사용할 수 있는 데이터라도 모델 학습에는 제한될 수 있습니다. 보안 시스템 안에서는 안전했던 데이터도 프롬프트 복사, 검색 결과 포함, 로그 저장 과정에서는 위험에 노출될 수 있습니다. 따라서 데이터 스튜어드십은 데이터가 어디에서 어떻게 이동하고 활용되는지 가시화하고, 이를 통제할 수 있는 규칙과 기술적 수단을 마련하는 데 초점을 둡니다.

가장 흔한 문제는 일상적인 업무 과정에서 발생합니다. 팀이 업무 효율을 높이기 위해 공개 도구에 기밀 자료를 입력하거나, 검색 계층이 오래된 정책 문서를 노출하거나, 비식별 처리된 데이터가 다른 자료와 결합되면서 재식별될 수 있습니다. 공급업체 시스템과 연동하는 과정에서 현업 부서가 인지하지 못한 방식으로 데이터 처리 방식이 바뀌는 경우도 있습니다. 개인정보 보호위원회의 지침 역시 재식별 위험과 개인의 통제권 약화를 설계 단계에서 고려해야 할 실질적 문제로 강조하고 있습니다.

왜 지금 중요한가

AI 도입이 확대되면서 민감 데이터가 이동할 수 있는 경로도 크게 늘어났습니다. 데이터는 이제 프롬프트, 검색 파이프라인, 미세조정 데이터셋, 평가셋, 텔레메트리, 로그 등 다양한 경로를 통해 활용됩니다. 그러나 많은 조직은 여전히 데이터가 기록 시스템 안에 머문다는 전제 아래 개인정보 보호를 관리하고 있습니다.

개인정보보호위원회의 지침은 데이터 관리가 특정 저장소나 시스템에 한정된 문제가 아니라, 데이터 수명 주기 전반에서 다뤄야 할 과제임을 분명히 합니다. 이에 따라 설계 단계부터 개인정보 보호를 반영하고, 데이터 최소화와 투명성, 거버넌스 체계를 AI 개발과 활용 전 과정에 선제적으로 적용해야 합니다.

2026년에 변화하는 것

2026년에는 데이터 스튜어드십이 개인정보 보호 규정 준수를 넘어, AI가 정보를 어떻게 사용하고 조직이 데이터 활용의 책임을 어떻게 입증할 것인지에 대한 더 넓은 관리 체계로 확장되고 있습니다.

- **데이터 관리의 절차화:** 데이터 처리 기준이 출시, 조달, 변경 관리 등 일상 업무 절차에 포함됩니다. 이를 통해 팀은 개인정보와 중요 정보 보호를 정해진 기준에 따라 신속하게 판단할 수 있습니다.
- **데이터 사용 기준 강화:** 동의, 이용자 기대, 데이터 주권 원칙이 AI 데이터 활용과 재사용, 공유 기준에 더 큰 영향을 미치고 있습니다.
- **입증 책임 강화:** 리더는 AI 수명 주기 전반에서 민감 데이터가 적절히 수집·활용·보호되었음을 입증할 수 있어야 합니다.

개인정보 보호와 데이터 활용은 반드시 상충하는 관계가 아닙니다.

데이터 스튜어드십을 잘 갖춘 조직은 최소한의 보호 조치로
위험을 감수하는 대신, 신뢰를 바탕으로 더 적극적인 데이터 공유를
이끌어냅니다. 이는 장기적으로 경쟁사가 따라오기 어려운 AI 역량과
신뢰 기반의 경쟁 우위로 이어질 수 있습니다.

한국의 상황

한국의 개인정보 규제 당국은 AI의 입력·출력 데이터 모두에
보호 의무가 적용될 수 있음을 강조하고 있으며, 생성형 AI
도구에 민감 정보를 입력할 때 발생할 수 있는 위험도 경고하고
있습니다. 국내 규제는 보호와 활용의 균형을 맞추기 위해 AI
전 생애주기별 개인정보 보호 기준과 가명정보 처리 방법을
구체화하며, 학습 데이터 활용을 둘러싼 불확실성을 줄여가고
있습니다.

이러한 흐름과 함께 정보주체의 권리를 강화하는 마이데이터
정책도 있습니다. 마이데이터는 개인이 자신의 데이터를 더 주
도적으로 관리하고 활용할 수 있도록 하는 방향으로 확대되고
있습니다. 동시에 감독 체계도 사후 제재 중심에서 사전 예방과
기술적 보호 조치 중심으로 이동하고 있습니다. 따라서 동의와
최소화 원칙을 일회성 절차가 아니라 지속적인 거버넌스로 내
재화하는 조직만이 AI 환경에서 신뢰를 경쟁력으로 전환할 수
있습니다.

리더십 우선순위

- **민감 데이터 활용 범위 식별:** AI가 민감 데이터를 다루는
주요 업무 흐름과 의사결정 지점을 먼저 파악해야 합니다.
- **보호 원칙의 업무 내재화:** 개인정보 보호, 동의, 목적 제한,
데이터 최소화 원칙이 실제 업무 절차 안에서 작동하도록
해야 합니다.
- **증명과 대응 체계 마련:** 데이터 처리 기준이 지켜지고 있음을
간단히 확인·보고할 수 있어야 하며, 오류나 유출 위험이 발
생했을 때 신속히 대응할 수 있는 절차를 갖춰야 합니다.

변화의 신호

민감 데이터가 승인된 경로 안에서만 활용되고, 편의를 이유로
비공식 도구에 입력되는 사례가 줄어듭니다. 팀은 어떤 데이터
를 어떤 AI 목적에 사용할 수 있는지 공통된 기준을 공유하며,
예외 상황도 일관되게 처리합니다. 리더는 주요 AI 업무 흐름에
서 데이터가 어떻게 수집·활용·보호되었는지 신속히 확인할
수 있고, 문제 발생 시 수정 조치가 일상적인 절차에 따라 이루어
집니다.

AI 리더의 사고방식

개인 and 고객이 조직에 데이터를 제공한다는 것은 단순히 정보
를 넘기는 것이 아니라, 신뢰를 맡기는 일입니다. 리더는 데이터
를 단순한 자산이 아니라 신뢰 관계의 기반으로 바라봐야 합니
다. 동의는 한 번의 체크박스로 끝나는 절차가 아니라, 데이터가
어떻게 사용되고 보호되는지 지속적으로 설명하고 관리하는 과
정입니다. 강력한 데이터 스튜어드십은 개인정보 보호를 부담
이 아니라 신뢰와 경쟁력을 만드는 기반으로 전환합니다.

AI 신뢰와 보증은 AI를 한 번 검토하는 절차가 아니라, 운영 중에도 계속 검증하고 통제하는 체계입니다. 신뢰는 출시 시점이 아니라, 실제 운영 과정에서 입증되어야 합니다.

AI가 고객 경험, 직원 업무, 규제 대상 의사결정에 깊이 활용될 수록 신뢰는 한 번 선언하고 끝낼 수 있는 것이 아닙니다. 경영진은 환경이 바뀌더라도 AI 기반 성과가 정확하고 공정하며, 법과 정책에 부합하게 유지되고 있다는 확신을 가져야 합니다.

AI 신뢰와 보증은 이러한 확신을 근거로 뒷받침하는 체계입니다. 시스템이 실제로 무엇을 수행했는지, 원래 어떤 기준에 따라 작동해야 했는지를 연결하고, 데이터 입력, 모델과 프롬프트 버전, 업무 흐름 단계, 최종 결과를 추적할 수 있게 합니다. 또한 내부 리스크 조직, 내부 감사, 제3자 검증 등 독립적인 확인이 가능하도록 필요한 근거와 환경을 마련합니다.

이 의제는 책임 있는 AI를 대체하는 것이 아니라 보완합니다. 책임 있는 AI가 원칙과 정책, 통제의 방향을 정한다면, 신뢰와 보증은 그 통제가 설계한 대로 작동하고 있으며 조직이 이를 검토 과정에서 입증할 수 있음을 지속적으로 확인합니다.

표준 측면에서도 AI 보증 체계는 점차 구체화되고 있습니다. ISO/IEC 42001과 같은 글로벌 표준은 AI 관리 시스템 안에서 조직의 역할과 책임을 정리하고 있으며, NIST AI 위험 관리 프레임워크는 문서화 자체보다 실제 통제 수단의 작동 여부를 강조합니다. 이는 AI 관련 표준이 단순한 가이드라인을 넘어, 기존의 데이터 보안과 규정 준수 체계와 함께 운영할 수 있는 관리 프레임워크로 성숙하고 있음을 보여줍니다.

왜 지금 중요한가

AI의 성과와 리스크는 고정되어 있지 않습니다. 모델 업데이트, 업무 흐름 변경, 데이터 패턴 변화, 자동화 확대는 시스템 장애가 발생하지 않더라도 AI 결과를 바꿀 수 있습니다. 전통적인 보증 방식은 정적인 프로세스와 주기적 점검을 전제로 하지만, AI는 계속 변화하는 운영 시스템에 가깝습니다.

따라서 신뢰를 유지하려면 근거가 지속적으로 축적되어야 합니다. 정상적인 업무 과정에서 감사 추적 기록, 검증 결과, 통제 신호를 함께 만들어낼 수 있는 조직은 돌발 리스크를 줄이면서도 더 빠르게 AI를 확장할 수 있습니다. 또한 리스크 조직, 고객, 규제 기관의 질문이 제기된 뒤 뒤늦게 신뢰를 회복하는 데 드는 시간과 비용도 줄일 수 있습니다.

2026년에 변화하는 것

2026년에는 AI 보증이 사후 점검에 머무르지 않고, 출시 전 검토와 운영 중 모니터링을 포괄하는 상시 관리 체계로 전환되고 있습니다.

- **상시 모니터링 체계:** 모니터링과 통제 신호가 정상 운영 과정에서 작동하면서, 드리프트, 정책 위반, 품질 저하를 조기에 감지할 수 있습니다. 이러한 문제는 예외적 사고가 아니라 일상적인 운영 관리의 일부로 다뤄집니다.
- **자동화된 근거 축적:** 의사결정 경로, 데이터 계보, 승인 내역, 성과 지표가 업무가 진행되는 과정에서 자동으로 기록됩니다. 이를 통해 보증에 필요한 근거가 별도 작업이 아니라 출시와 운영의 부산물로 생성됩니다.
- **독립 검증의 정례화:** 검증과 이의제기는 명확한 역할 분리 아래 반복 가능한 절차로 운영됩니다. 이를 통해 검토가 특정 개인이나 일회성 판단에 의존하지 않고, 지속적이고 확장 가능한 보증 체계로 자리 잡습니다.

중요한 것은 AI 시스템이 배포 당시 신뢰 여부가 아닙니다.
지금 이 순간 시스템이 내리는 결정을 신뢰할 수 있는지,
근거로 입증할 수 있는지가 중요합니다.

앰비언트 트러스트(Ambient Trust)는 시스템이 안정적으로
작동한다는 의미를 넘어, 왜 그렇게 작동하고 있는지를
필요할 때 설명할 수 있는 상태를 뜻합니다.

한국의 상황

한국의 규제 환경은 일회성 준수를 넘어, 외부 검증을 통해 AI 신뢰성을 지속적으로 입증하는 방향으로 전환되고 있습니다. 국가 차원의 신뢰성 검·인증 체계는 국제표준과 연계해 실제 시스템 성능을 점검하는 방식으로 고도화되고 있으며, 출시 전 자체 점검을 위한 자가 검증 수단도 마련되고 있습니다.

고위험 AI에 대한 영향평가 제도가 구체화되면서, 외부 기관의 인증과 검증은 공공 조달과 주요 서비스 도입 과정에서 점점 더 중요한 요건이 되고 있습니다. 이에 따라 국내 기업은 배포 시점의 안전성만이 아니라, 모델과 데이터가 변하는 운영 환경 속에서도 시스템의 신뢰성과 관련 근거를 지속적으로 입증해야 합니다.

리더십 우선순위

앰비언트 AI보증 체계로 전환하기 위한 핵심 우선순위는 다음 두 가지입니다.

- **운영 과정에서 근거 축적:** 데이터 계보, 승인 내역, 성과 지표가 구축과 운영 과정에서 자연스럽게 기록되도록 해야 합니다. 보증이 별도 작업이나 개인의 노력에 의존하지 않도록, 근거 생성 자체를 운영 절차에 포함해야 합니다.
- **독립적 검증 역량 확보:** 내부 감사, 리스크 관리, 전문가 검토 등 검증과 보증 활동이 명확한 독립성을 갖고 이루어져야 합니다. 동시에 AI 출시와 운영 속도에 맞춰 반복 가능한 검증 주기를 마련해야 합니다.

변화의 신호

성숙한 조직은 지난 분기 보고서가 아니라 현재의 운영 데이터를 바탕으로 AI 시스템을 신뢰할 수 있는지 판단할 수 있습니다. 보증 절차는 배포 속도에 맞춰 작동하며, 신뢰성을 해치지 않으면서도 지속적인 출시를 가능하게 합니다. 또한 민원이나 사고 발생 이후 문제를 확인하는 것이 아니라, 신뢰 저하의 신호를 사전에 감지하고 대응합니다. 이는 신뢰가 규제 대응 수단을 넘어, AI 확장을 뒷받침하는 운영 역량으로 자리 잡고 있음을 보여줍니다.

AI 리더의 사고방식

AI 리더는 이제 “이 시스템을 신뢰할 수 있는가”라는 질문을 넘어, “지금 이 순간 시스템이 신뢰할 만한 상태인지 확인할 수 있는가”를 물어야 합니다. 상시 보증 체계를 갖춘 리더는 속도와 안전을 양자택일의 문제로 보지 않습니다. 대신 지속적인 모니터링과 검증 체계를 통해 빠른 확장과 신뢰 확보를 함께 가능하게 만듭니다.

AI 감독은 형식적인 절차가 아니라 실제 의사결정에 개입할 수 있는 통제 체계여야 합니다. 이를 위해 사람의 판단이 필요한 영역, 권한 범위, 긴급 상황에서의 개입 기준을 사전에 명확히 정해야 합니다.

인간 중심의 통제는 그 어느 때보다 중요해지고 있습니다. AI 시스템이 사람의 세부 지시 없이도 계획하고, 추론하고, 결정하고, 실행하는 방향으로 발전하면서 인간의 통제와 기계의 자율성 사이의 경계가 빠르게 흐려지고 있기 때문입니다.

앞으로는 AI 기반 자율 시스템이 특정 사업 기능을 운영하고, 전략적 판단을 내리며, 사람이 따라가기 어려운 속도와 규모로 시장에 영향을 미치는 상황이 현실화될 수 있습니다. 예를 들어 AI가 디지털 자산을 거래하고, 새로운 시장 기회를 분석하며, 기업의 지분을 확보하고, 복잡한 투자 결정을 자율적으로 수행하는 시나리오를 생각해 볼 수 있습니다. 이 과정에서 잘못된 정보나 내부 정보를 활용한 거래가 발생하더라도, 인간 책임자가 누구인지 불분명하거나 개입이 늦어진다면 시장 신뢰는 빠르게 훼손될 수 있습니다. 이는 먼 미래의 상상이 아니라, 이미 등장한 기술 역량이 결합될 때 충분히 고려해야 할 위험입니다.

근본적인 질문은 분명합니다. AI 시스템이 사람보다 빠른 속도로 판단하고, 더 많은 정보를 처리하며, 더 넓은 영역에서 자율적으로 작동하게 될 때, 조직은 그 시스템에 대해 어떻게 의미 있는 인간의 통제를 유지할 것인가입니다.

왜 지금 중요한가

많은 조직은 기존의 감독 체계만으로 AI 시스템을 충분히 통제하기 어렵다는 점을 인식하고 있습니다. AI의 자율성이 높아지고, 에이전트가 여러 업무에서 판단과 실행에 관여하게 되면 사후 점검 중심의 감독 방식은 한계를 드러낼 수 있습니다.

따라서 인간의 권한과 책임은 시스템 설계 단계부터 명확히 반영되어야 합니다. 어떤 결정은 사람의 사전 승인이 필요한지, 어떤 결정은 사람의 감독 아래 AI에 위임할 수 있는지, 어떤 상황에서는 사람이 직접 개입해야 하는지를 구분해야 합니다.

인간 중심의 통제는 AI 확산을 늦추는 장치가 아니라, 중요한 의사결정에서 책임성과 신뢰를 유지하기 위한 운영 기준입니다.

2026년에 변화하는 것

인간의 통제는 형식적 승인 절차를 넘어, AI 시스템의 설계와 운영에 내재화되는 방향으로 발전하고 있습니다.

- **감독 역량 강화:** 사람은 AI의 판단 근거와 이상 신호를 이해하고, 언제 개입하거나 중단해야 하는지 판단할 수 있어야 합니다.
- **실시간 거버넌스:** AI 시스템이 설정된 경계에서 벗어나거나 예상과 다르게 작동할 때 즉시 알리고 대응할 수 있어야 합니다.
- **책임성의 내재화:** AI가 자율적으로 작동하더라도 최종 책임이 불분명해지지 않도록 설계, 배포, 운영 과정에 책임 기준을 포함해야 합니다.

시스템 아키텍처에 관리자의 권한을 내재화하는 것이 중요합니다.

자율 시스템이 기존 통제 체계보다 더 빠르게, 더 넓은 영역에서 작동할수록 인간의 통제는 단순히 AI의 작동을 지켜보는 역할에서 벗어나, AI가 작동할 수 있는 경계와 개입 기준을 정하는 역할로 바뀌어야 합니다.

한국의 상황

한국은 인공지능 기본법 전면 시행에 따라 고영향 AI 사업자에 대한 '인간의 관리·감독'을 법적 의무로 명문화했습니다.

사업자는 시스템 오작동 등 비상 상황이 발생했을 때 즉시 개입해 작동을 중단하거나 변경할 수 있는 실질적 통제 장치를 갖춰야 합니다. 또한 책임 회피를 막기 위해 감독 책임자를 외부에서 식별할 수 있도록 공개해야 합니다. 이는 AI 전 수명 주기에 걸쳐 최종 책임이 인간에게 귀속되어야 한다는 규제 의지를 반영합니다.

리더십 우선순위

- **시스템 설계에 인간 권한 반영:** 인간의 의사결정 권한을 AI 시스템 설계에 직접 반영하고, 자율 운영의 경계를 명확히 설정해야 합니다. 또한 AI의 결정 속도에 맞춰 사람이 개입하거나 결정을 재정의할 수 있는 권한을 유지해야 합니다.
- **감독 역량 강화:** 인력은 AI 시스템을 단순히 사용하는 수준을 넘어, 관리하고 통제하는 역량을 갖춰야 합니다. 이를 위해 경계 설정, 패턴 모니터링, 드리프트 감지, 임계치 접근 시 개입 절차를 운영할 수 있도록 교육과 프로세스에 투자해야 합니다.
- **추적 가능한 책임 체계 구축:** 모든 자율적 행동이 시스템 설계, 배포 기준, 운영 경계에 관한 인간의 결정과 연결되도록 해야 합니다. 그래야 AI가 자율적으로 작동하더라도 책임 주체를 명확히 확인할 수 있습니다.

변화의 신호

- **실행 가능한 감독 인터페이스:** 감독 시스템이 AI의 추론과 행동을 이해 가능한 정보로 제공해, 사람이 근거를 갖고 개입할 수 있습니다.
- **실질적인 개입 권한:** 인간은 형식적 검토를 넘어, 필요할 때 AI의 결정을 중단하거나 변경할 수 있는 권한을 가집니다.
- **유연한 운영 경계:** AI 역량이 발전하더라도 중요한 의사결정에 대한 인간의 통제력이 유지되도록 권한과 경계가 함께 조정됩니다.

AI 리더의 사고방식

관리자의 감독은 AI 시스템의 속도를 늦추기 위한 장치가 아니라, AI가 빨라질수록 통제력을 유지하기 위한 조건입니다.

가장 위험한 것은 실제 권한은 없으면서 인간이 통제하고 있다는 착각만 주는 형식적 감독입니다. 리더는 조직이 AI의 행동을 이해하고, 필요할 때 이의를 제기하며, 방향을 바꿀 수 있는 실질적 권한을 유지하고 있는지 점검해야 합니다.

전략적 미래 사고

**급변하는 AI 환경에서 리더는 가능성을 확장하는 동시에,
새로운 위험을 예측하고 통제할 수 있어야 합니다.**

AI는 기업의 공격 범위와 시스템 리스크를 넓히고 있으며, 오류가 빠르게 확산될 수 있는 새로운 운영 환경을 만들고 있습니다. 동시에 AI는 컴퓨팅 용량, 네트워크, 에너지, 물, 기술 인력 등 국가 인프라와도 밀접하게 연결됩니다.

따라서 AI 전략은 성능과 확장성만이 아니라 보안, 안전, 인프라 제약, 환경 영향까지 함께 고려해야 합니다. 변화 신호를 읽고 선택지를 유지하며, 필요할 때 방향을 전환할 수 있는 회복탄력성이 핵심입니다.



AI는 보안 위협의 범위를 프롬프트, 학습 데이터, 연결 시스템, 자율 에이전트까지 넓히고 있습니다. 이로 인해 대규모 조작과 데이터 유출, 사기 위험이 커지고 있습니다.

AI의 공격 범위는 기존 엔드포인트와 네트워크를 넘어 프롬프트, 학습·검색 데이터, 모델 공급망, 연결된 도구, 행동 권한을 가진 에이전트로 확장되고 있습니다. 그 결과 AI의 출력을 조작하거나, 정보를 빼내거나, 승인되지 않은 결과를 유도하는 방식의 공격이 더 다양해지고 있습니다. 이러한 공격은 기존 보안 경보 체계에서 쉽게 포착되지 않을 수 있습니다.

적대적 AI는 주요 위험 요소입니다. 프롬프트 인젝션, 데이터 오염, 모델 추출, 기만적 입력 등은 AI가 잘못된 결론을 내리거나 의도하지 않은 방향으로 작동하게 만들 수 있습니다. 공격자는 시스템이 생성한 행동이나 출력을 악용해 사람이 눈치채기 어려운 정보를 만들어내는 것을 목표로 합니다.

따라서 보안의 초점은 AI 시스템의 행동을 좌우하는 입력, 데이터, 권한, 연결 지점으로 확장되어야 합니다. AI는 반복 작업과 개인화된 공격의 비용을 낮추기 때문에 고객 사기, 임직원 사칭, 대규모 맞춤형 공격이 더 쉬워질 수 있습니다.

적대적 공격이 까다로운 이유는 공격자가 모델 자체뿐 아니라 주변 시스템까지 노릴 수 있기 때문입니다. 프롬프트는 침입 경로가 되고, 검색 커넥터는 데이터 유출 경로가 되며, 권한을 가진 에이전트는 공격자를 대신해 행동하는 실행자가 될 수 있습니다. 이제 에이전트의 실패는 단순한 오류가 아니라, 실제 보안 사고로 이어질 수 있는 위험입니다.

쉽게 말해 AI 보안은 이제 시스템 침입을 막는 것만으로 충분하지 않습니다. AI가 잘못된 행동을 하도록 유도되는 것을 막고, 문제가 발생했을 때 이를 빠르게 포착하며, 결과가 확산되기 전에 통제할 수 있어야 합니다.

왜 지금 중요한가

AI 보안은 이제 AI가 만들어내는 새로운 접점과 실행 권한까지 포함합니다. 프롬프트, 검색 커넥터, 학습 데이터, 모델 공급망, 행동 권한을 가진 에이전트가 모두 잠재적 공격 경로가 될 수 있습니다.

보안의 핵심은 AI가 생성한 결과를 단순히 보호하는 데 그치지 않습니다. 결과 조작, 민감 정보 유출, 승인되지 않은 행동, 허위 정보 생성까지 통제할 수 있어야 합니다. 많은 AI 실험이 충분한 통제 없이 운영되는 상황에서는 작은 취약점도 실제 업무와 고객 피해로 이어질 수 있습니다.

2026년에는 AI 보안이 모델 보안을 넘어 데이터 파이프라인, 에이전트 권한, 도구 연결, 실행 로그, 사고 대응까지 포함하는 운영 보안 체계로 확장될 것입니다.

2026년에 변화하는 것

보안은 이제 모델이 어떻게 호출되고, 어떤 데이터가 공급되며, 어떤 시스템과 연결되고, 어떤 권한으로 실행되는지를 전제로 설계됩니다.

- **공격 범위의 확장:** 프롬프트, 검색, 학습 데이터, 플러그인, 에이전트 권한, 도구 호출이 모두 보안 통제 대상이 됩니다.
- **대규모 조작의 용이화:** 사회공학 공격과 딥페이크, 가짜 신원, AI 기반 기만 전술이 더 낮은 비용으로 확산됩니다.
- **고위험 실행 주체가 되는 에이전트:** 비인가 행위자가 에이전트를 악용할 경우 신원, 권한, 로그, 신속한 차단 조치가 필요합니다.
- **적대적 테스트의 정례화:** 조작, 데이터 유출, 사기 시나리오를 포함한 공격 관점의 테스트가 상시 운영 절차로 자리 잡습니다.
- **AI의 방어적 활용:** AI는 위협 탐지와 대응 역량을 높이는 보안 도구로도 활용됩니다.

AI는 공격자가 하나의 취약점을 더 빠르고 다양한 방식으로 악용할 수 있게 만듭니다. 작은 약점도 기계적인 속도로 확산되어 큰 피해로 이어질 수 있습니다.

한국의 상황

한국도 글로벌 시장과 마찬가지로 AI 기반 사이버 위협의 영향을 받고 있습니다. 분산된 인력 구조와 AI 기반 업무 프로세스를 운영하는 국내 조직은 합성 사기와 사회공학 공격에 취약할 수 있습니다. 딥페이크와 가짜 신원을 활용한 개인화된 사기 기법은 더 정교해지고 있으며, 무심코 사용한 내부 AI 도구에서도 데이터 유출 위험이 발생할 수 있습니다.

이에 따라 국가 사이버 보안 지침과 역량 강화 프로그램도 AI 관점에서 재해석되고 있습니다. 이제 보안은 네트워크와 장비 보호에 그치지 않고, 모델, 프롬프트, 데이터 파이프라인, 에이전트 권한 설정까지 포함해야 합니다. 설계 단계부터 보안을 적용하는 것이 실질적인 보안 요건으로 자리 잡고 있습니다.

리더십 우선순위

- **기본 보안 체계 정비:** 강력한 디지털 신원 관리와 데이터 보호 체계를 갖춰야 합니다. AI는 접근 관리와 데이터 분류의 약점을 빠르게 드러냅니다.
- **AI 인터페이스 보호:** 프롬프트, 검색 커넥터, 도구, 에이전트 권한을 단순한 기능이 아니라 핵심 보안 통제 대상입니다.
- **공급망 보안 강화:** 모델 출처, 업데이트, 제3자 구성 요소, 데이터 입력에 대한 검증 체계를 세우고 신뢰성을 관리해야 합니다.
- **사이버 보안과 AI 통제의 통합:** 보안 조직, 대응 절차, 책임 기준을 정비해 AI 리스크가 핵심 사이버 리스크로 관리되도록 해야 합니다.
- **적대적 테스트의 상시 운영:** 조작, 데이터 유출, 사기 시나리오를 포함한 공격 관점의 테스트를 지속적으로 운영해야 합니다.

변화의 신호

AI 보안의 성숙도는 AI 위험을 핵심 사이버 리스크로 다루고, 이를 관리할 책임 주체와 상시 통제 수단을 갖추고 있는지에서 드러납니다. 문제가 발생했을 때는 어떤 정보가 노출됐고, 어떤 행동이 수행됐으며, 어떤 시스템에 영향을 미쳤는지 로그로 확인할 수 있어야 합니다.

또한 프롬프트와 검색 커넥터, 에이전트 권한이 모두 보안 아키텍처 안에서 관리되고, 주요 업무 흐름에 대해 공격 관점의 테스트와 개선 조치가 반복적으로 이뤄져야 합니다. 이때 조직의 AI 보안 역량은 체계적으로 성숙해집니다.

AI 리더의 사고방식

AI 보안 리더십은 시스템을 보호하는 데서 나아가, AI의 행동을 통제하는 방향으로 확장됩니다. 핵심은 권한과 접근이 집중되는 지점, 입력값이 조작될 수 있는 지점, 작은 취약점이 빠르게 확산될 수 있는 지점을 찾는 것입니다.

그리고 이러한 위험을 사후에 막는 것이 아니라, 설계 단계부터 차단할 수 있도록 보안 통제를 우선순위에 두어야 합니다.

안전 및 시스템 리스크

23

AI 안전은 모델과 에이전트, 업무 흐름이 서로 연결된 환경에서 운영 안정성을 유지하는 것을 의미합니다. AI의 자율성이 높아질수록 사전 검증과 대응 체계의 중요성도 커지고 있습니다.

보안이 공격자의 조작과 외부 위협을 다룬다면, 안전은 AI 시스템이 복잡한 업무 환경에서도 안정적으로 작동하는지를 다룹니다. 모델, 에이전트, 사람, 업무 흐름이 연결될수록 작은 오류도 다른 시스템으로 확산될 수 있습니다.

AI 안전의 핵심은 여러 구성 요소가 함께 작동할 때 발생할 수 있는 영향을 관리하는 데 있습니다. AI가 업무 프로세스 전반에 편입되면 작은 오류도 긴밀히 연결된 시스템을 따라 확산될 수 있습니다. 고위험 자율 기능은 특정 영역의 최적화가 전체 시스템의 취약점으로 이어질 수 있고, 합성 콘텐츠나 재활용 데이터가 반복되면 시간이 지날수록 품질 저하 문제가 발생할 수 있습니다.

이 때문에 AI 안전 업무는 점점 더 운영 관리에 가까워지고 있습니다. 의존성 파악, 연쇄적 실패 가능성 점검, 강력한 비상 대응 절차, 고위험 자율성에 대한 안전장치가 함께 요구됩니다. 안전한 AI 개발은 사후에 검토하는 항목이 아니라, 시스템 수명 주기 전반에 포함되어야 하는 운영 기준입니다.

대부분의 리스크가 그렇듯, 안전과 시스템 리스크의 이해관계자는 사이버 보안 팀에만 국한되지 않습니다. 운영 리더십, 전사 리스크, 기술, 법무, 감사, 도메인 책임자가 함께 업무 결과를 이해하고 책임 기준을 정해야 합니다. 이들이 시스템 의존성에 대한 공통된 이해를 갖고, 예상과 다르게 움직일 때 신속히 조정할 수 있어야 AI 안전 관리가 제대로 작동합니다.

왜 지금 중요한가

AI는 점점 더 긴밀하게 연결된 시스템으로 바뀌고 있습니다. 모델과 에이전트, 업무 흐름이 계속 상호작용하면서 서로의 행동에 영향을 미치기 때문입니다. 따라서 AI 안전의 초점은 개별 모델의 성능을 넘어, 스트레스 상황에서도 시스템이 안정적으로 작동하는지 확인하는 방향으로 이동하고 있습니다.

2026년에 변화하는 것

AI가 운영 전반에 깊이 편입되면서, 안전은 개별 모델의 문제가 아니라 업무 흐름 전체의 안정성과 회복탄력성, 제한된 자율성을 다루는 시스템 차원의 관리 과제가 되고 있습니다.

- **의존성 가시화:** 조직은 모델과 에이전트, 도구, 제3자가 어디에서 어떻게 연결되는지 파악해야 합니다. 이를 통해 리더는 실패가 연쇄적으로 확산될 수 있는 지점을 확인할 수 있습니다.
- **안전성 검증의 일상화:** 고위험 자율 기능은 출시 전 검토에 그치지 않고, 완화 조치와 판단 근거, 에스컬레이션 경로를 포함한 안전성 검증이 지속적으로 요구됩니다.
- **비상 대응 체계 마련:** AI 시스템에 문제가 발생했을 때 중단, 기능 제한, 사람에게 인계하는 절차가 표준 운영 방식으로 마련되어야 합니다.
- **운영 안정성 사전 점검:** 시나리오 리허설과 회복탄력성 훈련, 기본 제어 장치의 반복 검증을 통해 AI 시스템의 준비 상태를 확인하는 과정이 자리 잡고 있습니다.

AI가 운영에 도입되면 작은 오류도 특정 영역의 문제로 끝나지 않습니다. 연결된 시스템을 따라 확산되며, 사람이 개입하기 전에 더 큰 문제로 변질 수 있습니다.

한국의 상황

한국의 AI 안전 체계는 산업 진흥과 신뢰 기반 조성을 함께 추진하는 방향으로 빠르게 제도화되고 있습니다. AI가 공공서비스와 고위험 영역, 높은 신뢰가 요구되는 업무에 활용되면서 효과성과 안전성을 함께 관리해야 한다는 기대가 높아지고 있습니다.

인공지능기본법 시행, AI 안전연구소 출범, 공공부문 초거대 AI 도입 가이드라인, 생성형 AI 개인정보 처리 안내서 등은 한국에서도 AI 안전 리스크와 운영 안정성이 더 이상 개별 기술 문제가 아니라 거버넌스와 운영의 핵심 과제가 되고 있음을 보여줍니다.

리더십 우선순위

- **시스템 의존성 파악:** 모델과 에이전트, 도구, 업무 흐름이 어떻게 연결되어 있는지 파악해야 합니다. 이를 통해 리더는 장애가 연쇄적으로 확산될 수 있는 지점을 미리 확인할 수 있습니다.
- **비상 대응 체계 마련:** 중요한 자율 기능은 제한된 범위 안에서 작동해야 하며, 문제가 발생했을 때 사람이 개입하거나 시스템을 중단할 수 있는 절차를 갖춰야 합니다.
- **안전성 검증 체계 확보:** 고위험 자율 기능에 대해서는 위험 요소와 완화 조치, 판단 근거, 에스컬레이션 경로를 포함한 검증 체계를 마련해야 합니다.
- **이해관계자 조율:** 운영, 리스크, 기술, 법무, 감사, 도메인 책임자가 함께 대응할 수 있는 협업 체계를 구축해야 합니다. 안전은 단순한 기술 문제가 아니라 전사적 운영 과제로 다뤄져야 합니다.

변화의 신호

- **의존성 파악:** 리더는 중요한 AI 기반 프로세스가 어떤 시스템과 업무 흐름에 의존하는지 명확히 파악할 수 있습니다.
- **운영 경계 설정:** 고위험 자율 기능은 명확한 비상 대응 절차, 중단 장치, 복구 경로와 함께 운영됩니다.
- **신속한 대응 체계:** 대응 기준과 에스컬레이션 경로가 명확해 부서 간 협업이 빠르게 이뤄질 수 있습니다.

AI 리더의 사고방식

핵심은 안정성입니다. 리더는 개별 모델의 성능보다 여러 구성 요소가 함께 작동할 때 운영 시스템이 회복탄력성을 유지할 수 있는지 확인해야 합니다. 신뢰는 정상 상황에 대한 기대가 아니라, 스트레스 상황에서도 시스템이 견딜 수 있다는 근거에서 나옵니다.

한국의 AI 경쟁력은 데이터센터와 클라우드 인프라, 네트워크 연결성 인프라가 얼마나 안정적으로 뒷받침되느냐에 달려 있습니다. 여기에 AI 시스템을 구축·운영할 수 있는 기술 인력과 국방·보건 등 핵심 영역의 신중한 기술 선택이 더해져야 합니다.

AI 역량은 소프트웨어만으로 결정되지 않습니다.

데이터센터, 로컬 컴퓨팅과 클라우드 리전, 네트워크 용량, 전력과 냉각, 숙련된 인력은 AI를 안정적으로 운영하기 위한 핵심 기반입니다. 이러한 기반이 충분히 갖춰지지 않으면 상용 AI 서비스와 운영의 회복탄력성을 유지하기 어렵습니다.

이러한 인프라 관점은 AI 도입이 개별 기업의 선택을 넘어 국가 역량과 직접 연결된다는 점을 보여줍니다. 국내 AI 생태계의 성과와 품질은 네트워크가 분산 운영을 뒷받침할 수 있는지, 전력·냉각 인프라가 제약 없이 관리되는지, 필요한 기술 인력이 적기에 확보되는지에 따라 달라집니다. 특히 공공서비스, 국방, 보건처럼 국가적 영향력이 큰 영역에서는 기술 선택과 인프라 운영에 대한 신중한 판단이 필요합니다.

리더에게 중요한 변화는 AI 전략 수립에서 생태계 협력이 필수 과제가 되고 있다는 점입니다. 조직은 표준화, 기술 개발, 보증 기준 형성에 참여해야 하며, 데이터센터 준비성, 전력·냉각 인프라, 보안 클라우드 가용성, 회복탄력성 있는 거버넌스 체계를 갖추기 위해 정부, 연구기관, 산업계와 더 적극적으로 협력해야 합니다.

왜 지금 중요한가

AI는 더 이상 기업 내부의 기술 도입 문제에 머물지 않습니다. 데이터센터, 전력, 냉각, 기술 인력, 네트워크, 클라우드 인프라가 함께 갖춰져야 AI의 확산 속도와 운영 안정성을 뒷받침할 수 있습니다. 앞으로 한국의 AI 도입 속도는 이러한 인프라와 정책 환경이 얼마나 함께 발전하느냐에 따라 달라질 것입니다.

2026년에 변화하는 것

AI 역량은 점점 더 국가 인프라의 선택과 제약에 영향을 받게 됩니다. 데이터센터, 전력, 기술 인력, 클라우드 인프라가 기업의 AI 전략을 둘러싼 중요한 환경 요소가 되고 있습니다.

- **데이터센터의 전략적 중요성 확대:** 데이터센터의 용량과 입지, 회복탄력성은 대규모 AI 활용이 얼마나 실현 가능하고 경제적이며 신뢰할 수 있는지를 결정합니다.
- **핵심 인프라 선택의 중요성 확대:** 주요 AI 업무를 어떤 인프라에서 운영할지, 어떤 클라우드와 컴퓨팅 자원을 활용할지에 대한 판단이 리스크와 안정성, 공공 신뢰와 연결됩니다.
- **기술 인력 확보의 중요성 증가:** AI 시스템을 구축하고 운영하는 역량은 모델이나 플랫폼 접근성만큼 중요한 경쟁 요소가 되고 있습니다.
- **생태계 협력 확대:** 조직은 AI 기반 인프라를 함께 형성하는 표준, 파트너십, 국가 프로그램에 더 적극적으로 참여해야 합니다.

AI는 국가 역량을 좌우하는 핵심 요소로 자리 잡고 있습니다. 데이터센터와 에너지, 냉각 인프라, 기술 인력, 네트워크 자원의 제약이 커질 경우 국가 AI 전략의 추진 속도도 둔화될 수 있습니다.

한국의 상황

최근 국내 AI 산업은 데이터센터의 수도권 집중으로 인한 송전망 한계와 전력계통 부담이라는 물리적 제약에 직면하고 있습니다. 이에 따라 AI 역량 확보의 핵심은 단순한 모델 접근성을 넘어, 전력 수급과 계통 연계 가능성을 확보하는 방향으로 옮겨가고 있습니다.

정부도 전력계통영향평가를 통해 수도권 집중을 관리하고, 지역 분산과 전력망 확충을 위한 인프라 개선을 추진하고 있습니다. 결과적으로 송전망은 단순 설비가 아니라 국가 AI 역량을 좌우하는 핵심 기반으로 부상했으며, 데이터센터 입지와 전력 확보 능력은 기업의 규모 확장, 비용 구조, 장기 회복탄력성을 결정하는 전략 과제가 되고 있습니다.

리더십 우선순위

- **생태계 참여:** 정부와 연구기관, 산업계와 협력해 표준, 기술 인력 기반, 보증 기준을 함께 만들어가야 합니다.
- **제약을 고려한 계획:** 데이터센터와 에너지, 물, 네트워크가 규모와 비용, 회복탄력성에 미치는 영향을 전략적으로 요소로 다뤄야 합니다.
- **핵심 영역의 인프라 선택:** 국내 통제와 안정성이 중요한 영역을 정의하고, 아키텍처와 조달 기준을 그에 맞게 설계해야 합니다.

변화의 신호

국가 인프라 관점이 기업 전략에 어떻게 반영되는지는 조직의 상황과 규모에 따라 다르게 나타날 수 있습니다. 중요한 첫 단계는 데이터센터, 전력, 냉각, 네트워크, 기술 인력과 같은 인프라 제약이 AI 전략에 어떤 영향을 미치는지 구체적으로 논의하는 것입니다.

AI 리더의 사고방식

리더는 데이터센터와 네트워크, 에너지, 기술 인력, 핵심 영역의 인프라 선택을 하나의 포트폴리오 관점에서 바라봐야 합니다. 이러한 요소는 평상시에는 잘 드러나지 않지만, 위기 상황이나 수요가 급증하는 시점에 더 큰 영향을 미칠 수 있습니다. 따라서 회복탄력성과 선택권을 중시하고, 장기적으로 운영 여지를 넓히는 설계와 파트너십을 마련하는 것이 중요합니다.

제로 이미션 인텔리전스(Zero Emission Intelligence)는 AI의 탄소 발자국과 에너지 · 물 사용량을 관리하고, AI 운영이 ESG 목표와 함께 작동하도록 만드는 접근입니다.

제로 이미션 인텔리전스는 AI를 폭넓은 ESG 체계 안에서 바라봅니다. 에너지 수요와 수자원 사용, 공급망 영향, 시스템 배포와 운영 방식이 가져오는 사회적 결과까지 함께 고려합니다.

핵심은 AI가 본질적으로 유해하다는 의미가 아니라, AI가 빠르게 확장되는 기술이라는 점입니다. AI가 반복 업무와 고객 접점, 제품 기능에 깊이 편입될수록 환경적 영향도 함께 커질 수 있습니다.

데이터센터는 AI 운영의 보이지 않는 비용 구조를 드러내는 중요한 지표입니다. AI는 전력과 물을 연산 역량으로 전환하며, 규모가 커질수록 이러한 투입 요소는 비용과 리스크, 운영 승인에 직접적인 영향을 미칩니다. AI 도입에 따른 환경 부담은 데이터센터를 통해 다음과 같이 확인할 수 있습니다.

- **전력 수요의 급증:** 국제에너지기구는 전 세계 데이터센터 전력 사용량이 AI 연산 수요 증가에 따라 빠르게 늘어날 것으로 전망하고 있습니다.
- **전력 점유율 확대:** 데이터센터가 세계 전력 소비에서 차지하는 비중은 현재보다 더 커질 것으로 예상됩니다.
- **상당한 수자원 소비:** 대규모 데이터센터는 냉각을 위해 많은 양의 물을 사용할 수 있으며, 물 사용량 역시 주요 관리 대상이 되고 있습니다.

운영 측면에서 효율성은 ESG와 비용을 함께 관리하는 핵심 수단입니다. 모델 선정, 라우팅 전략, 프롬프트 설계, 업무 흐름 설계, 스케줄링, 인프라 배치는 모두 환경 발자국에 영향을 미칩니다. AI 요구 수준이 높아질수록 워크로드도 주요 배출원과 수자원 소비원으로 인식되고 있습니다. 따라서 지속적으로 측정하고 설명할 수 있으며, 개선 가능한 관리 대상으로 다뤄야 합니다.

왜 지금 중요한가

AI의 환경적 영향은 지속가능성 리더만의 과제가 아닙니다. 재무, 조달, 운영, 리스크 담당자에게도 중요한 경영 이슈가 되고 있습니다. 기술 성과만 보고 AI를 확장하면 비용, 전력 사용, 물 사용, 규제와 평판 리스크를 함께 키울 수 있습니다.

따라서 기업은 AI 사용 규모와 환경 영향을 초기 단계부터 함께 관리해야 합니다. AI의 영향은 일상적인 사용량이 쌓이면서 커지며, 비용과 배출, 사업 운영의 정당성에도 영향을 미칠 수 있습니다.

2026년에 변화하는 것

ESG 기대가 운영 단계에서 AI와 더 밀접하게 결합되면서, 효율성은 책임 있는 AI 확장과 환경 발자국 관리의 핵심 운영 기준이 되고 있습니다.

- **환경 영향의 지표화:** 탄소, 에너지, 물 사용량은 AI 운영 과정에서 추적 · 관리해야 할 핵심 지표로 자리 잡습니다.
- **효율성 중심의 운영 설계:** 모델 규모, 라우팅, 업무 흐름 설계는 성능을 유지하면서 자원 사용을 줄이는 방향으로 조정됩니다.
- **조달과 보고 체계의 고도화:** 인프라와 공급업체 선택은 지속가능성, 거버넌스, 공시 요구를 함께 고려하는 기준으로 평가됩니다.
- **AI 기반 감축 확대:** AI 투자는 운영과 공급망 전반에서 측정 가능한 배출 감축 성과와 연결됩니다.

모든 AI에는 실제 자원이 투입됩니다.

AI의 성능 뒤에는 전력과 물, 탄소, 인적 노동이라는 보이지 않는 비용이 존재합니다.

한국의 상황

국내 AI 산업의 확장은 탄소 회계와 ESG 공시 측면에서도 중요한 과제가 되고 있습니다. 그동안 많은 기업이 AI를 클라우드 기반 서비스로 인식해 왔지만, 실제로는 데이터센터 전력 수요와 냉각, 공급망, 운영 방식이 함께 영향을 미칩니다.

상대적으로 제한적인 재생에너지 접근성, 무탄소 전력 확보의 어려움, 전력망 제약은 향후 AI 확장의 중요한 조건이 될 수 있습니다. 또한 냉각 과정에서 발생하는 물 사용과 전력 효율도 데이터센터 입지와 운영 전략의 주요 판단 기준으로 부상하고 있습니다.

결국 워크로드 배치 최적화, 재생에너지 조달, 탄소를 고려한 데이터 운영, 공급업체 선택은 비용 절감과 탄소 감축을 동시에 달성하기 위한 조건입니다. 나아가 기업의 AI 사업이 장기적으로 정당성을 확보하기 위한 필수 요건이 될 것입니다.

리더십 우선순위

제로 이미션 인텔리전스는 AI의 탄소 발자국을 단순한 기술 문제가 아니라 비용, 운영 회복탄력성, 평판 전반에 영향을 미치는 경영 과제로 다룹니다. 따라서 리더십은 AI 확장 과정에서 환경 영향을 측정하고 관리할 수 있는 기준을 마련해야 합니다.

이를 위해 모델 효율, 워크로드 배치, 공급업체 기준, 보고 체계에 관한 선택을 관행이나 기본값에 맡기지 않아야 합니다. AI가 거버넌스와 운영의 일부로 성숙할수록 논의의 범위는 탄소를 넘어 에너지, 물, 공급망 영향, 사회적 결과까지 포함하는 ESG 의제로 확장됩니다.

변화의 신호

AI는 ESG 관리 체계 안으로 편입되고 있으며, AI의 탄소 발자국은 비용과 품질, 리스크와 함께 관리 가능한 성과 지표로 다뤄지고 있습니다. 에너지와 탄소, 물 사용량은 명확한 책임 소재 아래 주요 AI 워크로드 중심으로 추적됩니다.

또한 투자와 배치 결정은 단순한 성능 향상이나 비용 절감만이 아니라, ESG 요건에 부합하는 인프라와 운영 방식으로 이어지고 있습니다.

AI 리더의 사고방식

리더가 AI를 중요한 기업 인프라로 이해할 때, AI의 탄소 발자국도 개선 가능한 성과 지표로 관리할 수 있습니다.

미래예측 거버넌스는 시나리오와 변화 신호를 바탕으로 작은 규모의 선택지를 유지하고, 특정 기술이나 공급업체에 대한 과도한 의존을 피하며, 변화가 본격화될 때 신속하게 대응할 수 있도록 준비하는 접근입니다.

미래예측 거버넌스는 현재의 계획을 넘어 미래의 변화를 미리 살피고, 주요 가정을 검증하며, 변화가 본격화되기 전에 선택지를 마련하는 관리 체계입니다. AI 환경에서 이 접근이 중요한 이유는 변화의 속도가 일정하지 않기 때문입니다. 모델 출시와 도구 혁신, 규제 변화, 컴퓨팅 비용은 사업 환경을 단기간에 크게 바꿀 수 있습니다.

기업에 필요한 것은 미래를 정확히 맞히는 능력이 아니라, 여러 가능성에 대비하는 운영 방식입니다. 효과적인 미래예측 거버넌스는 의사결정에 의미 있는 시나리오를 만들고, 변화의 방향을 보여주는 조기 신호를 포착하며, 상황에 따라 확대하거나 중단할 수 있는 선택지를 남겨둡니다. 이를 통해 특정 기술이나 공급업체에 과도하게 의존하는 것을 피하고, 전환이 필요한 시점에 빠르게 움직일 수 있습니다.

AI는 리스크와 생산성, 신뢰를 동시에 재편하는 유동적인 영역이기 때문에 거버넌스도 함께 진화해야 합니다. 리더십은 외부 환경 변화에 맞춰 전략적 가정을 재검토하고, 자율성에 대한 위험 수준을 조정하며, 아키텍처와 파트너십, 통제 수단을 바꿀 수 있는 운영 리듬을 갖춰야 합니다.

핵심은 불확실한 미래를 하나로 단정하지 않는 것입니다. 미래예측 거버넌스는 미래를 정확히 예측하기 위한 활동이 아니라, 조직이 변화의 신호를 빠르게 읽고 필요할 때 방향을 전환할 수 있는 탄력성을 확보하는 방법입니다.

왜 지금 중요한가

2025년의 변화는 AI 환경이 짧은 기간 안에도 크게 바뀔 수 있음을 보여줬습니다. 오픈 웨이트 모델의 부상, 에이전트 생태계의 변화, 추론 기능의 확산, 규제 환경의 변화는 기존 전략 계획만으로 대응하기 어려운 속도로 진행되고 있습니다.

따라서 AI 리더는 중요한 선택을 미루는 것이 아니라, 조기 신호를 활용해 의사결정의 가역성을 유지해야 합니다. 기술과 파트너십을 유연하게 설계하고, 다음 변화에 대비하는 태도가 더 중요한 역량으로 떠오르고 있습니다.

2026년에 변화하는 것

전략은 AI 역량, 규제, 경제 환경이 1년 안에도 실질적으로 달라질 수 있다는 점을 반영해 더 짧은 주기와 적응형 거버넌스로 이동합니다.

- **전략 검토 주기의 단축:** AI 전략을 연간 계획에 묶어두지 않고, 시장 변화와 기술 신호에 맞춰 더 자주 점검하고 조정합니다.
- **변화 신호의 체계적 관리:** 핵심 신호를 선별해 경영진의 의사결정에 반영하고, 전환 시점을 더 빠르게 포착합니다.
- **선택지를 남기는 투자:** 특정 기술이나 공급업체에 과도하게 종속되지 않도록 소규모 실험과 대안 옵션을 유지합니다.
- **거버넌스의 상시 조정:** 위험 수준, 자율성 설정, 통제 수단을 환경 변화에 맞춰 지속적으로 조정하는 체계로 전환합니다.

최근의 AI 변화는 중요한 전환점이 매년 나타날 수 있음을 보여줍니다.

장기 전략도 변화 신호를 읽고 빠르게 조정할 수 있는 체계를 갖추지 못하면 금세 유효성을 잃을 수 있습니다.

한국의 상황

한국 기업은 글로벌 AI 역량을 빠르게 도입하면서도 위험 관리와 규제 준수를 중시하고 있습니다. 이에 따라 AI는 일회성 과제가 아니라 상시적으로 관리해야 하는 전략 의제로 다뤄지고 있습니다.

정부와 공공연구기관의 미래 예측, 기술 정책 논의, 산업별 AI 적용 전략은 기업의 전략 수립에도 중요한 참고 기준이 되고 있습니다. 다만 기존의 연간 계획 주기만으로는 시장 변화와 조기 신호에 충분히 대응하기 어렵습니다. 산업별 활용 방식과 기술 성숙도, 규제 환경이 빠르게 달라지는 만큼, 전략적 가정과 위험 신호, 역량 배분을 수시로 점검하는 체계가 필요합니다.

리더십 우선순위

- **전략 주기 단축:** AI 전략을 연간 계획에 묶어두지 않고, 시장 변화에 맞춰 수시로 점검하고 조정합니다.
- **시나리오 실행 기준 마련:** 주요 시나리오별 변화 신호와 대응 기준을 정해, 필요한 때 바로 실행할 수 있도록 합니다.
- **경영진 거버넌스 강화:** AI를 단순한 기술 도입 과제가 아니라 자율성, 회복탄력성, 외부 의존성을 함께 관리하는 상시 거버넌스 의제로 다룹니다.

변화의 신호

- **전략 검토 리듬:** 전략은 연간 계획 주기에만 의존하지 않고, 실제 환경 변화와 시장 신호에 맞춰 더 빠른 주기로 검토됩니다.
- **선택지 확보:** 조직은 기술, 규제, 경제 환경이 달라질 때 확대하거나 중단할 수 있는 대안을 보유합니다.

AI 리더의 사고방식

AI 리더는 현재의 사업을 운영하는 동시에 미래의 사업을 준비하는 데 익숙해져야 합니다. 앞서가는 리더는 전략을 고정된 계획이 아니라 살아 있는 시스템으로 운영합니다.

핵심은 불확실성을 없애는 것이 아니라, 변화가 나타났을 때 빠르게 움직일 수 있는 선택권을 남겨두는 것입니다. 2026년의 경쟁력은 AI 거버넌스를 일회성 계획이 아니라, 변화 신호에 맞춰 전략과 운영을 조정하는 체계로 만드는 데서 나옵니다.

다음 단계의 리더십



다음 단계로의 도약

AI 네이티브 기업의 다음 단계는 기술 도입을 넘어 경영진의 전략적 판단으로 옮겨가고 있습니다. 향후 12~24개월 동안 AI가 손익 구조를 어떻게 바꿀지, 자본과 컴퓨팅 자원을 어느 수준까지 투입할지, 신뢰와 안전의 기준을 어디에 둘지, 통제력을 유지하면서 어떤 파트너와 협력할지 결정해야 합니다.

한국 기업은 제조 기반, 풍부한 운영 데이터, 높은 디지털 수용성을 강점으로 갖고 있습니다. 다만 이러한 강점이 실제 경쟁력으로 이어지려면 리더십이 AI를 기술 과제가 아니라 사업 성과와 운영 방식의 문제로 다뤄야 합니다. AI 확장의 속도와 방향은 최고경영진의 판단과 실행 체계에 따라 결정됩니다.

AI, 비즈니스 운영의 중심으로

2026년 AI는 주변부 혁신 과제를 넘어 기업 운영의 중심으로 이동하고 있습니다. 다음 단계의 AI는 아키텍처와 운영 체계, 의사결정 방식, 리스크 관리, 고객 접점과 서비스 방식 전반에 영향을 미칩니다.

AI가 만들어내는 변화는 더 빠른 실행, 높은 정확성, 단위 경제성 개선으로 나타납니다. 동시에 투명성, 도입 수준의 불균형, 컴퓨팅 비용 증가, 운영 과정에서의 신뢰 확보라는 과제도 커지고 있습니다. 리더십의 역할은 이러한 기회와 위험을 함께 읽고, 전략적 판단을 통해 AI가 만들어낼 미래를 설계하는 데 있습니다.

최고 경영진의 결정

AI 운영의 방향과 속도는 몇 가지 핵심 선택에 달려 있습니다. 리더십은 AI가 손익 구조를 어떻게 바꿀지 분명히 이해해야 합니다. AI 관련 지출은 단순한 기술 예산이 아니라, 사업 운영 방식과 미래 경쟁력에 대한 전략적 투자로 봐야 합니다.

자본과 컴퓨팅 용량은 의식적인 투자 영역이 됩니다. 필요한 용량을 확보하고, 성과를 내는 분야에 자원을 배분하며, 어떤 데이터와 인프라 파트너를 선택할지 결정해야 합니다. 출시 기준에는 출처 관리, 개인정보 보호, 영향 수준에 따른 설명 가능성, 안전성 검토가 포함되어야 하며, 문제가 발생했을 때 적용할 명확한 대응 절차도 필요합니다.

파트너십은 역량과 통제력을 기준으로 선택해야 합니다. 어떤 모델과 공급업체에 의존할지, 지식재산과 데이터는 어떻게 보호할지, 협력 구조가 장기적으로 기업에 어떤 영향을 미칠지 따져봐야 합니다. 이 모든 판단은 AI 확장의 속도를 늦추기 위한 장치가 아니라, 확장을 지속 가능하게 만드는 경영 기준입니다.

AI 확장을 위한 핵심 역량

AI 실행은 반복 가능한 성과로 이어질 때 의미가 있습니다. 이를 위해 네 가지 역량이 함께 필요합니다

1. 역량 (AI 이해)

AI가 사업의 경제성을 어디서, 어떻게 바꾸는지에 대한 공통의 이해가 필요합니다. 경영진은 AI를 단순한 생산성 도구가 아니라 사업 구조를 바꾸는 요소로 이해해야 합니다. 구성원은 AI의 활용 범위와 한계, 필요한 검토 기준을 알아야 합니다. 명확한 기대를 바탕으로 속도와 품질의 균형을 맞출 때 AI 활용이 실제 성과로 이어질 수 있습니다.

2. 전략 (가치 발견)

AI가 가치를 창출할 수 있는 영역을 식별하고, 이를 사업 우선순위와 투자 판단으로 연결해야 합니다. 리더십은 고객 경험, 운영 효율, 리스크 관리, 신사업 기회 중 어디에 집중할지 명확히 정해야 합니다. 핵심은 모든 영역에 동시에 투자하는 것이 아니라, 성과 가능성이 높은 영역을 선별해 확장 가능한 구조로 설계하는 것입니다.

3. 파운드리 (가치 증명 및 확장)

조직에서 찾은 AI 기회를 반복 가능한 결과로 전환하려면 검증과 확장의 체계가 필요합니다. 리더는 투자수익률과 리스크에 대한 의사결정 기준을 두고, 재사용 가능한 표준 패턴과 운영 방식을 구축해야 합니다. 검증된 성과는 빠르게 확산하고, 효과가 낮은 과제는 조기에 조정할 수 있어야 합니다.

4. 신뢰 (가치 보호)

AI 확장은 신뢰를 전제로 해야 합니다. 책임 있는 AI, 거버넌스, 투명성은 별도 항목이 아니라 AI 운영 방식의 일부가 되어야 합니다. 사용 기준, 품질 검증, 감사 추적, 보안 점검, 성능 모니터링이 함께 작동할 때 AI는 안정적으로 확장될 수 있습니다. 신뢰는 규제를 피하기 위한 장치가 아니라, AI가 사업 운영 방식에 자리 잡게 만드는 조건입니다.

성숙한 AI 리더십은 무모한 확장과 다릅니다.

학습 체계에 대한 지속적 투자와 명확한 윤리 기준, 불확실성을 인정하는 태도, 위험 관리와 실행 속도의 균형이 필요합니다.

주의 사항

AI 확장은 빠르게 추진하되, 통제 가능한 구조 안에서 이뤄져야 합니다. 파일럿이 여러 방향으로 분산되거나, 사업 성과보다 벤치마크와 기술 지표에 치우치거나, AI 도구가 기준 없이 확산되면 운영 리스크가 커질 수 있습니다. 중요한 것은 도입 속도를 높이는 것만이 아니라, 실행 속도를 뒷받침할 거버넌스와 운영 기준을 함께 갖추는 것입니다.

컴퓨팅과 데이터 사용량이 늘수록 비용 관리의 중요성도 커집니다. 모델 사용량, 처리 비용, 데이터 활용 수준을 단위 경제성 관점에서 추적하면 비용 구조를 더 명확히 파악하고 개선 방향을 찾을 수 있습니다.

시스템 위험은 단일 모델의 실패를 넘어섭니다. 의존성이 복잡해지고, 규제·법적·사회적 영향이 함께 커지면 작은 오류도 큰 문제로 번질 수 있습니다. 예외 상황에 대응할 수 있는 상시 역량을 갖추고, 위험을 줄이는 구조를 사전에 설계해야 합니다.

경영진의 기준은 실용적이고 책임감 있어야 합니다. 성과를 입증하고, 효과가 있는 것은 확장하되, 그렇지 않은 것은 중단할 수 있어야 합니다. 학습 속도를 높이면서도 통제 기준을 유지하는 것이 AI 시대 리더십의 핵심입니다.

불확실성 속에서의 리더십

빠르게 변하는 AI 환경에서 리더십은 막연한 낙관이 아니라, 선택지를 열어두고 변화에 대응할 수 있는 운영 능력에서 나옵니다. 미래를 정확히 예측하는 것보다 중요한 것은 변화 신호를 빠르게 읽고, 필요할 때 전략을 조정할 수 있는 체계를 갖추는 것입니다.

강조점은 개별 모델의 성능 평가에서 벗어나고 있습니다. 이제는 데이터, 모델, 업무 흐름, 통제 장치가 함께 작동하는 기업 AI 시스템을 관리하는 역량이 중요합니다. 이 과정에서 성능, 속도, 비용, 책임성이 함께 관리되어야 합니다.

전체 시스템을 관리하는 역량은 기업의 중요한 경쟁력이 됩니다. 테스트와 시뮬레이션을 통해 AI 시스템의 영향을 사전에 점검하고, 실제 업무에 적용하기 전에 위험을 줄일 수 있어야 합니다. 기존 시스템과 충돌하지 않으면서도 조직의 운영 방식에 자연스럽게 결합되는 AI 구조가 필요합니다.

한국 기업은 제조 기반, 운영 데이터, 높은 디지털 수용성이라는 강점을 갖고 있습니다. 특히 복잡한 업무 흐름과 현장 데이터를 보유한 산업에서는 AI 성과를 빠르게 검증하고 개선할 수 있는 여지가 큼니다. 다만 이러한 강점은 자동으로 경쟁력이 되지 않습니다. 데이터 품질, 시스템 연계, 보안 기준, 책임 있는 활용 원칙이 함께 갖춰져야 합니다.

사전 예방 원칙은 계속 적용됩니다. 할 수 있는 것에 따라 움직이기보다, 시스템을 적용 가능한 방식으로 유지하고 문제가 커지기 전에 통제할 수 있어야 합니다. AI 리더십은 기술 확장 자체가 아니라, 그 확장이 기업의 운영 방식 안에서 책임 있게 자리 잡도록 만드는 데 있습니다.

AI는 더 이상 선택적 도입 과제가 아니라, 기업 경쟁력의 기본 조건입니다.

AI 도입은 더 이상 선택 사항이 아닙니다. 비용 절감과 생산성 향상, 변화 대응력을 확보하기 위한 기본 조건이 되고 있습니다.

그러나 시장을 선도하는 기업은 단순한 도입을 넘어, AI를 업무 방식과 의사결정 구조 전반에 반영하고 있습니다.

아이디어를 기업 가치로 연결하기 위해서는 빠르게 실험하고, 검증된 성과를 확장할 수 있는 실행 체계가 필요합니다.

PwC의 통합 접근법

아이디어에서 기업 가치 창출까지, 실행 속도와 확장 가능성을 함께 높입니다.



역량 Fluency

경영진과 구성원이 AI를 이해하고 실제 업무에 활용할 수 있도록 조직 전반의 AI 이해도와 실행 기반을 강화합니다.



전략 Strategy

AI 도입 목표를 구체화하고, 이를 사업 성과와 연결합니다. 우선순위를 정하고 전사 확장을 위한 실행 방향을 수립합니다.



파운드리 Foundry

AI의 가치를 검증하고, 검증된 솔루션을 확장 가능한 방식으로 구축해 실제 사업 성과로 연결합니다.



신뢰 Trust

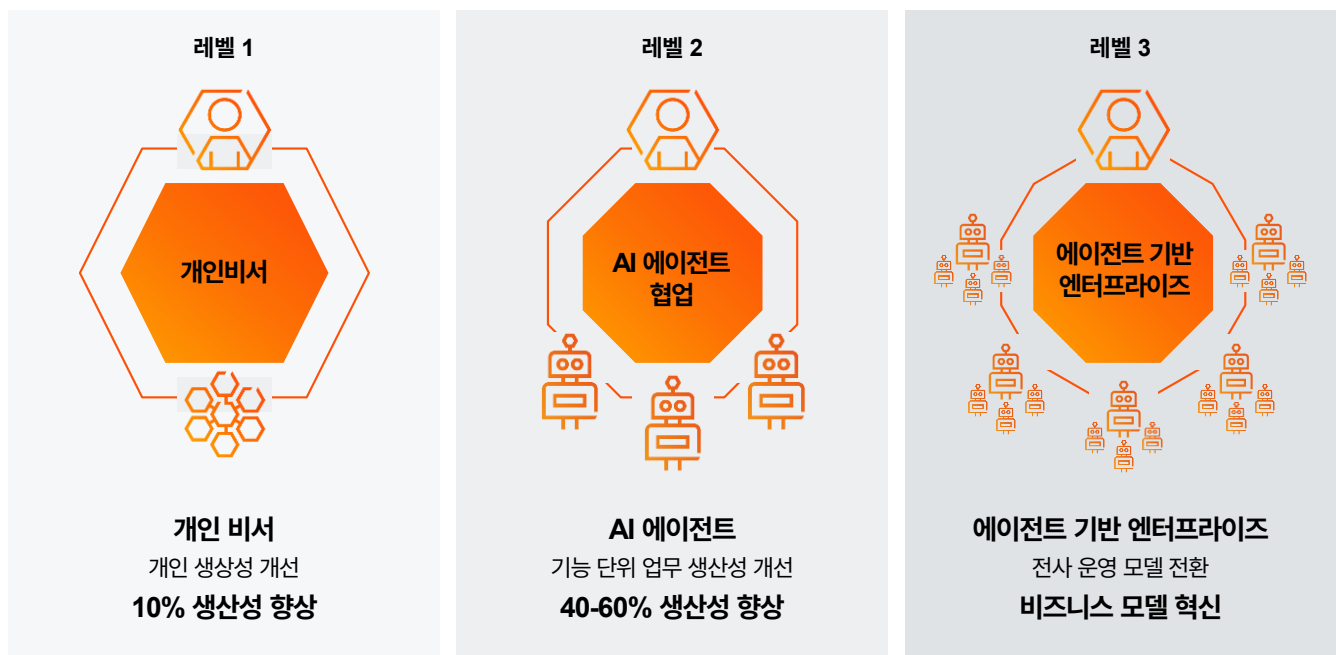
책임 있는 AI, 거버넌스, 투명성을 운영 체계에 반영해 신뢰할 수 있는 AI 성장 기반을 마련합니다.

역량, 전략, 파운드리, 신뢰가 연결될 때 AI는 단발적 파일럿을 넘어 실제 사업 성과로 확장될 수 있습니다. 조직은 검증된 과제를 더 빠르게 실행하고, 반복 가능한 방식으로 확산하며, AI 전환에 대한 신뢰를 높일 수 있습니다.

AI 활용 단계

사람의 역량을 확장하는 AI

기업의 AI 활용은 세 단계로 발전합니다. 첫 단계는 개인 업무를 지원하는 AI, 두 번째는 팀 업무를 보강하는 AI 에이전트, 세 번째는 조직 전반의 운영 방식을 바꾸는 AI입니다. 각 단계는 이전 단계를 기반으로 발전하며, 개인 생산성 향상에서 업무 흐름의 전면 자동화로 확장됩니다.



레벨 1. 개인 비서

개인 비서 단계의 AI는 초안 작성, 요약, 분석, 검색 등 일상 업무에서 개인의 생산성을 높이는 도구입니다. 이 단계의 AI는 업무 구조를 바꾸기보다, 개인이 맡은 일을 더 빠르고 효율적으로 처리하도록 돕는 데 초점이 있습니다.

예시: 컨설턴트가 AI 비서를 활용해 이메일 초안을 작성하고, 회의를 요약하며, 초기 슬라이드를 만드는 방식입니다. 이를 통해 개인 업무 처리 속도는 높아지지만, 팀 전체의 업무 흐름이 구조적으로 바뀌는 것은 아닙니다.

레벨 2. AI 에이전트 협업

AI 에이전트 협업 단계에서는 AI가 팀 프로세스에 포함되어 일부 업무를 수행하고, 사람이 이를 검토하거나 승인합니다. 이 단계의 AI는 특정 업무 방식을 재구성해 수작업을 줄이고 일관성을 높입니다.

예시: 감사 팀이 AI 에이전트를 활용해 작업 문서 초안을 준비하고, 규정 준수 여부를 확인하며, 증거 자료를 분류할 수 있습니다. 이때 직원은 판단이 필요한 영역, 리스크 검토, 고객과의 대화에 더 집중하게 됩니다.

레벨 3. 엔터프라이즈 자동화

에이전트 기반 엔터프라이즈 단계에서는 AI 에이전트가 개별 업무 보조를 넘어 여러 부서와 핵심 시스템을 연결하는 전사 업무 흐름에 편입됩니다. 이 단계에서는 AI를 전제로 업무 처리 방식, 역할 분담, 의사결정 절차가 다시 설계됩니다.

예시: 청구 처리 전 과정이 자동화될 수 있습니다. 한 에이전트는 문서를 수집하고, 다른 에이전트는 자격 요건을 확인하며, 또 다른 에이전트는 의사결정 내용을 정리하고 고객 커뮤니케이션 초안을 작성합니다. 사람은 복잡한 판단이나 예외 처리가 필요한 사안을 검토하고 최종 결정을 내립니다.

1. 벤치마크 재설정: 업무 성과로 이동하는 모델 평가 기준

2025년 AI 모델 평가는 단순한 지능 비교에서 실제 업무 수행 능력을 확인하는 방향으로 더 뚜렷하게 이동했습니다. 업계의 관심은 모델이 얼마나 높은 점수를 받는지뿐 아니라, 제한된 조건과 실제 업무 맥락 안에서 얼마나 안정적으로 성과를 내는지로 확장되고 있습니다.

이에 따라 추론, 코딩, 전문 지식, 장기 과제 수행, 에이전트 실행 능력을 평가하는 벤치마크의 중요성이 커졌습니다.

평가 기준도 한층 까다로워졌습니다. 고난도 전문 지식과 복합 추론 능력을 평가하는 벤치마크가 등장하면서, AI가 단순히 정답을 맞히는 수준을 넘어 복잡한 문제를 어느 정도까지 처리할 수 있는지 확인하려는 시도가 이어졌습니다. 이는 모델 평가의 초점이 평균적인 응답 품질에서 실제 업무에 가까운 문제 해결 능력으로 이동하고 있음을 보여줍니다.

성과도 일부 영역에서 뚜렷하게 나타났습니다. 주요 모델은 고난도 수학 추론 평가와 코딩 평가에서 높은 성과를 보였고, 일부 에이전트형 코딩 시스템은 실제 GitHub 이슈를 해결하는 방식으로 개발 업무 지원 가능성을 입증했습니다. 과거에는 전문가의 판단이 필요하다고 여겨졌던 영역에서도 AI가 의미 있는 성과를 보이기 시작한 것입니다.

모델 제품군도 세분화되고 있습니다. 일상 업무를 위한 빠른 범용 모델, 복잡한 분석을 위한 추론 특화 모델, 소프트웨어 개발을 지원하는 에이전트형 코딩 모델, 대규모 입력을 처리하는 긴 컨텍스트 모델, 텍스트 · 이미지 · 오디오 · 비디오를 함께 다루는 멀티모달 모델이 각기 다른 역할을 맡고 있습니다. 2025년 말 기준 AI는 단순한 초안 작성과 요약 넘버, 대규모 정보를 처리하고 여러 단계의 업무 흐름을 지원하는 방향으로 발전했습니다.

2. AI의 비용 구조: 컴퓨팅에서 탄소까지

2025년의 중요한 변화 중 하나는 AI가 더 이상 가볍게 확장되는 소프트웨어가 아니라는 점이 분명해졌다는 것입니다. AI의 경제성은 모델 성능만으로 판단하기 어렵습니다. 인프라 비용, 컴퓨팅 자원, 추론 비용, 에너지 사용량까지 함께 관리해야 하는 운영 과제가 되고 있습니다.

최고 수준의 모델 학습에는 막대한 비용이 들고, 추론 비용 역시 지속적인 운영비로 자리 잡고 있습니다. 대형 모델을 활용해 고급 업무를 처리하려면 상당한 GPU 자원이 필요하며, 프롬프트 처리, 데이터 이동, 스토리지, 재시도, 보안 · 안전 검토까지 포함하면 실제 비용은 단순한 호출 비용보다 훨씬 커질 수 있습니다.

국내에서는 이러한 비용 문제가 데이터센터 전력 수요와 직접 연결되고 있습니다. AI 데이터센터는 고성능 서버와 냉각 설비를 필요로 하기 때문에 전력 사용량이 크고, 입지 선정에서도 전력 확보가 핵심 조건으로 부상하고 있습니다. 특히 수도권에 데이터센터와 전력 수요가 집중되면서 전력망 제약과 지역 분산 필요성이 함께 제기되고 있으며, AI 사용이 대규모로 반복될수록 탄소 배출과 냉각 비용 부담도 커질 수 있습니다.

따라서 AI 비용 관리는 단순한 IT 비용 관리가 아닙니다. 모델 선택, 작업 배분, 인프라 배치, 에너지 사용, 탄소 영향까지 함께 고려하는 경영 과제가 되고 있습니다. 2026년에는 AI 확장과 함께 비용 효율성과 환경 영향을 동시에 관리하는 체계가 더 중요해질 것입니다.

3. 추론의 실행 기능화: 생각하는 모델에서 검증하는 시스템으로 전환

2025년 AI의 또 다른 전환점은 추론 기능이 고급 모델의 품질 요소에 머물지 않고, 명시적으로 호출하고 관리할 수 있는 기능으로 바뀌었다는 점입니다. AI는 더 이상 빠른 응답만 제공하는 시스템이 아니라, 근거를 확인하고 계산하며 검증 가능한 답을 내도록 설계되는 방향으로 이동하고 있습니다.

이 변화는 업무 적용 방식에도 영향을 미쳤습니다. 중요한 것은 모델이 언제 깊이 생각해야 하는지, 언제 정보를 찾아야 하는지, 언제 계산과 검증을 수행해야 하는지를 구분하는 것입니다. 추론은 모델 내부의 능력에 그치지 않고, 업무 설계와 검토 절차의 일부가 되고 있습니다.

2026년에는 추론 자체가 거버넌스의 대상이 될 것입니다. 깊은 사고에는 시간과 비용이 들기 때문에, 모든 업무에 같은 수준의 추론을 적용하기보다 가치가 큰 영역을 선별해야 합니다. 조직은 유창한 답변보다 신뢰할 수 있는 결과에 비용을 지불하게 될 것이며, 빠른 응답이 필요한 업무와 깊은 추론이 필요한 업무를 구분하는 능력이 중요해질 것입니다.

4. 에이전트의 현실화: 자율성보다 통제·권한· 모니터링

에이전트는 2025년에 크게 주목받았지만, 초기의 과장된 기대와 실제 적용 사이에는 차이가 있었습니다. 기술적으로는 계속 발전했지만, 자율성이 커질수록 권한 관리와 실행 통제의 중요성도 함께 커졌습니다.

기업이 에이전트를 실제 업무에 투입하려면 단순히 기능을 구현하는 것만으로는 충분하지 않습니다. 권한 부여, 실행 기록, 감사 추적, 도구 접근, 실패 시 복구가 함께 설계되어야 합니다. 에이전트가 어떤 시스템에 접근하고, 어떤 작업을 수행했으며, 문제가 발생했을 때 누가 개입할 수 있는지 명확해야 합니다.

2026년의 에이전트는 자율성 자체보다 통제 가능한 실행 구조를 중심으로 발전할 것입니다. ID 관리, 최소 권한 기반의 도구 접근, 실시간 중단·재개 기능, 실행 기록, 업무 상태 확인 체계가 중요해집니다. 이러한 기반이 갖춰질 때 에이전트는 실험적 자동화 도구를 넘어 실제 배포 가능한 업무 시스템으로 발전할 수 있습니다.

5. 오픈 웨이트 모델의 확산: 단일 모델에서 포트폴리오 전략으로

2025년에는 특정 공급업체의 폐쇄형 모델에만 의존하지 않아도 충분한 성능을 확보할 수 있다는 인식이 확산되었습니다. 오픈 웨이트 모델은 학습 방식이 모두 공개되는 것은 아니지만, 공개된 가중치를 바탕으로 사용자가 모델을 실행하고 조정할 수 있다는 점에서 선택지를 넓혔습니다.

DeepSeek의 R1 모델과 주요 소형 변형 모델, Meta의 Llama 4 추진은 다운로드 가능한 모델이 실험 단계를 넘어 실제 운영 환경에서도 활용될 수 있음을 보여줬습니다. 일부 조직은 오픈 웨이트 모델을 활용해 비용을 낮추고, 특정 업무에 맞춘 모델 운영 가능성을 검토하기 시작했습니다.

이에 따라 AI 운영 방식은 하나의 최적 모델을 찾는 방식에서 여러 모델을 조합해 관리하는 방식으로 바뀌고 있습니다. 다운로드 가능한 오픈 웨이트 모델은 공급업체 API 의존도를 낮추고, 조직이 업무 특성에 맞게 모델을 배치할 수 있게 합니다.

경쟁은 모델 자체보다 운영 체계의 문제로 이동하고 있습니다. 기업은 에이전트를 위한 인증과 접근 통제, 모델 간 연동 방식, 오픈 에이전트 프로토콜, 도구 상호운용성 등을 함께 고려해야 합니다. 앞으로의 AI 도입은 모델 하나를 선택하는 문제가 아니라, 목적에 맞는 모델 포트폴리오를 구성하고 관리하는 역량에 달려 있습니다.

6. 멀티모달 AI의 확산: 콘텐츠 생성에서 출처 확인까지

텍스트, 이미지, 오디오, 비디오를 함께 해석하고 생성하는 멀티모달 AI는 2025년에 실용적 가능성을 크게 넓혔습니다. 비디오 생성은 더 이상 실험적 기술에 머물지 않고 제품화 단계에 들어섰고, 이미지 생성, 오디오 동기화, 대화형 음성 효과, 출처 확인 도구도 함께 발전했습니다.

이 변화는 콘텐츠 제작 방식뿐 아니라 콘텐츠의 신뢰 문제를 함께 부각시켰습니다. 디지털 콘텐츠가 대량으로 생성될수록 진위 확인, 출처 추적, 탐지 기술이 더 중요해집니다. AI가 만든 콘텐츠가 빠르게 확산되는 환경에서는 생성 능력만큼이나 출처와 사용 이력을 확인할 수 있는 체계가 필요합니다.

2026년에는 콘텐츠 전반에 걸쳐 출처 확인과 추적 기능이 기본 인프라로 자리 잡을 것입니다. 합성 미디어 탐지, 워터마킹, 출처 표시, 생성 이력 관리가 별도 기능이 아니라 시스템 설계 단계에서 함께 고려되어야 합니다. 멀티모달 AI의 경쟁력은 생성 품질뿐 아니라 신뢰할 수 있는 유통과 검증 체계를 갖췄는지에 따라 달라질 것입니다.

7. 규제 대응의 내재화: 배포 가능한 AI의 조건

2025년에는 규제와 지침이 단순한 법무 검토를 넘어 AI 아키텍처와 배포 방식에 영향을 미치기 시작했습니다. 국내에서는 인공지능기본법을 중심으로 AI의 신뢰성, 투명성, 안전성에 대한 요구가 커졌고, 규제 대응은 배포 가능한 AI 시스템의 필수 조건이 되고 있습니다.

이제 기업은 추적 가능성과 검증 체계를 갖춘 문서화, 프리프로덕션 모니터링, 도구 사용 기록, 합성 미디어 책임 체계를 마련해야 합니다. AI 시스템은 개발 이후 규제를 맞추는 방식이 아니라, 설계 단계부터 로그, 평가, 출처 관리, 안전 모니터링을 포함해야 합니다.

2026년의 핵심은 실행 가능한 거버넌스입니다. 거버넌스는 배포를 늦추는 통제 장치가 아니라, AI를 안전하고 빠르게 확장하기 위한 운영 인프라가 되어야 합니다. 배포 가능한 AI는 기술 성능뿐 아니라 검증, 기록, 모니터링, 책임 기준을 함께 갖춘 시스템이어야 합니다.

2026년의 기술 트렌드

2025년의 전환점은 2026년 AI 경쟁의 기준이 독립형 모델이 아니라 AI 시스템으로 이동하고 있음을 보여줍니다. 핵심은 업무에 맞는 모델 조합, 빠른 응답과 깊은 추론의 구분, 도구 오케스트레이션, 통제 가능한 에이전트, 검증을 내장한 거버넌스입니다.

2026년의 전략은 추론의 경제성, 에이전트 제어 체계의 성숙도, 오픈 웨이트와 폐쇄형 모델을 함께 활용하는 포트폴리오 관리, 출처 확인과 모니터링을 기본으로 포함하는 AI 시스템 구축에 집중될 것입니다.

용어집

에이전트 / 에이전트 AI: 인간의 감독 하에 단계적 작업을 수행하고, 도구 및 데이터와 상호작용하며, 워크플로우 업무를 완수할 수 있는 AI 시스템.

에이전트 운영체제(AOS): 일관된 안전성, 규정 준수 및 자원 통제를 위해 에이전트 군집을 실행하고 관리하는 표준화된 계층.

AGI(범용 인공지능): 단일 AI가 인간이 할 수 있는 모든 지적 작업을 이해하고 학습하며 수행할 수 있는 이론적인 미래 역량.

AI 리스크: 부정확성, 편향성, 오용, 데이터 유출 및 통제력 상실을 포함하여 AI 사용으로 인해 발생하는 운영, 윤리, 법률 및 평판 관련 위험.

AI 신뢰: AI 시스템이 규제 및 조직의 기대에 부합하여 안전하고 정확하며 일관되게 작동한다는 확신.

편향성: 편향된 데이터, 비대표적인 학습 또는 결함이 있는 가정으로 인해 AI 결과물에서 발생하는 체계적 오류.

사고의 연쇄(Chain-of-Thought) 프롬프팅: 명확성, 정확성 및 신뢰성을 향상하기 위해 AI에게 추론 단계를 보여달라고 요청하는 프롬프팅 기법.

컴퓨테이션(연산): AI 시스템을 학습하고 실행하는 데 필요한 처리 능력으로, 종종 규모, 비용 및 속도에 대한 제약 요인이 됨.

데이터 주권: AI와 데이터를 신뢰할 수 있는 지역 및 직접 통제하는 제공업체 내에 유지하는 것.

디지털 트윈: 물리적 자산, 프로세스 또는 시스템을 시뮬레이션하고 결과를 예측하는 데 사용되는 고충실도 가상 복제본.

드리프트(성능 저하): 실제 환경이나 데이터의 변화로 인해 시간이 지남에 따라 AI 모델의 성능이 감소하는 현상.

엣지 AI: 더 낮은 지연 시간을 위해 로컬 기기(휴대전화, 노트북, NPU)에서 효율적인 모델을 실행하는 것.

설명 가능성: AI 시스템이 특정 결과물이나 결정을 내린 이유를 이해할 수 있는 능력.

파운데이션 모델: 광범위한 데이터를 사전 학습하여 다양한 특정 작업에 맞게 조정(파인튜닝)할 수 있는 대규모 AI 모델. 파운데이션 모델의 예로는 OpenAI의 GPT, Google의 Gemini, Meta의 Llama, Anthropic의 Claude 등이 있음.

생성형 AI: 학습된 패턴을 사용하여 텍스트, 코드, 이미지 또는 요약과 같은 새로운 콘텐츠를 생성할 수 있는 AI.

생성형 워크플로우: AI가 초안 작성, 변환 및 단계적 작업을 처리하고 인간은 판단과 예외 사항에 집중하도록 재설계된 프로세스.

GPT(생성형 사전 학습 트랜스포머): 트랜스포머 아키텍처를 사용하여 구축되고 인간과 유사한 텍스트를 생성하도록 학습된 대규모 언어 모델 제품군.

가드레일: AI를 안전하고 규정을 준수하며 정확하고 브랜드 이미지에 부합하게 유지하는 규칙, 통제 및 점검 사항.

환각(Hallucination): AI가 부정확하거나 조작되었거나 실제 데이터에 근거하지 않은 정보를 확신을 가지고 생성하는 현상.

인간 중심의 의사결정(Humans at the Helm): AI를 인간의 판단을 대체하는 것이 아니라 보완하는 용도로 사용하며, 주요 결정에 대해 인간이 책임을 지도록 보장하는 거버넌스 원칙.

거대언어모델(LLM): 방대한 텍스트 데이터셋으로 학습되어 인간과 유사한 언어를 이해하고 요약하며 생성할 수 있는 AI 모델.

머신러닝: 명시적으로 프로그래밍되지 않아도 데이터에서 패턴을 학습하여 예측이나 결정을 내리는 AI의 한 방식.

멀티모달리티: 텍스트, 이미지, 오디오, 비디오 등 두 가지 이상의 데이터 유형을 처리하고 생성할 수 있는 AI의 능력.

자연어 처리(NLP): AI가 인간의 언어를 이해하고 다룰 수 있게 하는 기술.

신경망: 상호 연결된 계층(뉴런)을 사용하여 데이터에서 복잡한 패턴을 학습하는, 뇌의 구조에서 영감을 받은 모델.

관측 가능성(Observability): AI 시스템이 어떻게 행동하고 수행하는지에 대한 엔드투엔드 가시성(에이전트가 무엇을 했는지, 어떤 데이터를 다루었는지, 어떤 결과물을 생성했는지 등).

오케스트레이션: 모델, 도구 및 단계를 하나의 원활한 프로세스로 결합하는 것.

패턴: AI 워크플로우, 에이전트 및 통제 장치를 구축하기 위한 재사용 가능한 청사진(예: 작업 라우팅, 가드레일 적용 또는 에이전트 모니터링 방법).

파일럿에서 프로덕션으로: 소규모 실험에서 비즈니스 영향력을 갖춘 배포 및 확장된 시스템으로 전환하는 것.

프롬프팅: AI 모델에 지시를 내리는 행위(정확하고 신뢰할 수 있는 결과물을 생성하기 위해 프롬프트를 작성하는 방법 포함).

RAG(검색 증강 생성): 사실적 정확성을 향상하기 위해 신뢰할 수 있는 내부 또는 외부 데이터 소스에 AI 응답의 근거를 마련하는 방법.

강화 학습: 올바른 행동에는 보상을 주고 잘못된 행동에는 불이익을 주는 피드백을 통해 AI가 학습하는 기법.

책임 있는 AI(Responsible AI): AI가 공정하고 윤리적이며 안전하고 약속에 부합하도록 보장하기 위한 정책 및 관행.

지속 가능한 AI: 에너지 효율성과 환경적 영향을 최소화하도록 AI를 설계하는 것.

합성 데이터: 민감한 정보를 노출하지 않으면서 실제 데이터를 모방하여 실험 환경 구축을 지원하는 AI 생성 데이터셋.

텔레메트리(원격 측정): AI의 성능, 가치, 사용량, 정확성, 비용 및 위험을 보여주는 실시간 데이터 신호.

학습(Training): AI 모델이 대량의 고품질 데이터에 노출되도록 하여 특정 작업을 수행하고 패턴을 인식하며 결정을 내리도록 가르치는 것.



본 보고서 소개

본 간행물은 인공지능이 2026년과 그 이후 기업에 미칠 영향에 대한 PwC의 분석과 전망을 제시합니다.

PwC의 연구와 기업 협업 경험을 바탕으로, 주요 산업 동향과 새롭게 부상하는 AI 활용 방식을 함께 다룹니다.

본 보고서는 ISO/IEC 42001, NIST AI 위험관리 프레임워크 등 국제 표준과 국내 인공지능기본법 및 관련 정책·지침을 참고했습니다. 외부 표준과 기술 동향은 맥락 파악을 위해 활용했으며, 분석과 해석, 결론은 PwC의 견해입니다.

방법론 및 미래 예측

본 보고서는 글로벌 AI 트렌드를 종합하되, 한국 시장 상황과 규제 환경, 비즈니스 맥락을 함께 고려했습니다. 미래 예측은 현재 확인되는 기술 발전과 시장 변화에 근거합니다.

다만 AI의 발전 속도가 빠른 만큼, 변화가 예상과 다른 방향으로 전개될 가능성도 있습니다. 따라서 본 보고서의 인사이트는 확정적 예측이 아니라, 전략적 방향을 검토하기 위한 참고 자료로 활용해야 합니다.

인사이트와 사례 연구, 전문 지식을 제공해 주신 한국 및 글로벌 PwC 팀원들께 감사드립니다. 또한 AI가 기업에 미치는 실질적 영향에 대한 당사의 이해를 높이는 데 도움을 주신 고객사와 여러 기관에도 감사드립니다.

AI 활용에 대한 고지

본 보고서 작성 과정에서 생성형 AI 도구를 책임 있고 투명하게 활용했습니다. 생성형 AI는 트렌드 스캐닝, 연구 내용 종합, 콘텐츠 구조화 등을 보조하는 역할로 사용했습니다.

AI가 생성하거나 제안한 내용은 모두 전문가가 검토하고 필요한 내용을 반영했습니다. 이러한 접근은 투명성, 인간의 통제, 책임성을 중시하는 책임 있는 AI 원칙을 반영합니다.

중요 참고 사항

본 간행물은 법률 또는 규제 자문을 목적으로 하지 않습니다. 보고서의 내용을 실제 의사결정에 활용할 때에는 관련 규제와 조직의 구체적 상황을 별도로 검토해야 합니다.

출처

본 보고서에서 분석한 AI 환경을 형성하는 데 기여한 연구자, 기관 및 실무자들의 노고에 깊은 감사를 표합니다. 당사의 사고에 특히 많은 통찰을 제공한 대표적인 출처는 다음과 같습니다:

- PwC 글로벌 CEO 설문조사 및 “제29차 글로벌 CEO 설문조사”
- PwC 글로벌 AI 일자리 바로미터
- PwC 디지털 신뢰 인사이트 설문조사
- PwC 글로벌 사례 연구 및 연구 자료

표준 및 프레임워크

- ISO/IEC 42001
- ISO/IEC 23894
- ISO/IEC 38507
- NIST AI 위험관리 프레임워크

한국 규제·정책 관련 출처

- 국가법령정보센터(AI 기본법)
- AI 기본법 시행령 입법예고 / AI 규제합리화 로드맵
- 개인정보보호위원회_생성형·공개된 개인정보 처리 안내서
- 개인정보보호법·행정기본법·자동화된 결정 조치기준
- 과학기술정보통신부 (국가 AI 윤리기준·신뢰할 수 있는 AI 실현전략)
- 금융위원회 (금융분야 AI 활용 보도자료)
- 국가인공지능전략위원회·국가인권위원회 (AI 인권영향평가)
- 인공지능안전연구소(AISI)·International AI Safety Report 2025
- TTA_신뢰할 수 있는 AI 개발 안내서, AI 신뢰성 검·인증(CAT)

기술 및 산업 관련 출처

- 국제에너지기구(IEA) 데이터 센터 전력 소비 전망
- 세계경제포럼(WEF) 데이터 센터 물 소비 추정치
- 제11차 전력수급기본계획(KESIS)
- KDI 경제정보센터
- UNEP 데이터 센터 지속가능성 가이드
- 에너지경제연구원 세계에너지시장 인사이트
- 전기신문 「AI 인프라 3대 강국 과제」
- KM저널 「2025 국내 AI 모델 총정리」
- IEA AI 에너지 및 자원 소비 보고서
- Anthropic의 헌법적 AI(Constitutional AI) 연구
- Google DeepMind AI 안전 연구
- Meta AI 연구 (Llama 개발)
- NVIDIA AI 엔터프라이즈 연구
- Harvard Business Review AI 전략 연구
- 세계경제포럼(WEF) 미래 일자리 보고서

AI 연구, 워킹 그룹 및 벤치마크

- 케임브리지 대학교 미래 지능 센터: AI 미래 및 책임 연구
- MIT 정보 시스템 연구 센터(CISR)
- 스탠포드 인간 중심 인공지능 지수(HAI), 파운데이션 모델 투명성 지수(FMTI)
- 옥스퍼드 인터넷 연구소 AI 거버넌스 연구
- Humanity's Last Exam 벤치마크
- AIME 2025 수학적 추론 벤치마크
- 리눅스 재단 데이터 및 AI 연구
- AI 안전 및 정렬(Alignment)에 관한 학술 연구
- 세계경제포럼(WEF) AI 거버넌스 연합
- 유네스코(UNESCO) AI 교육
- 한국고용정보원
- 대한상공회의소

AX Node

이승환 Partner | AX Node Leader

seung-whan.lee@pwc.com
02-3781-9863

전용욱 Partner | AX Node Deputy Leader

yong-wook.jun@pwc.com
02-709-7982

Primary Contact

김재동 Partner

jae-dong.kim@pwc.com
02-3781-3467

이창훈 Partner

chang-hoon.lee@pwc.com
02-3781-1755

조홍래 Partner

hong-rae.cho@pwc.com
02-709-8434

정형근 Partner

hyung-geun.jung@pwc.com
02-709-3992

이권일 Director

kweon-il.lee@pwc.com
02-3781-9657

Domain Experts

나국현 Partner

gookhyun.na@pwc.com
02-3781-9119

이민지 Partner

min-ji.lee@pwc.com
02-3781-9200

김세준 Partner

sejoon.s.kim@pwc.com
02-2628-7384

정해민 Partner

hai-min.jeong@pwc.com
02-3781-0108

변장원 Partner

byun.jang-won@pwc.com
02-709-4764

이승욱 Partner

seung-wook.lee@pwc.com
02-709-7012

서종혁 Partner

jong-hyuk.seo@pwc.com
02-3781-1429

정수정 Partner

soo.jung.jeong@pwc.com
02-709-7038

조성재 Partner

sung-jae.jo@pwc.com
02-3781-9728

신민섭 Partner

min-sup.shin@pwc.com
02-709-8247



삼일회계법인

삼일회계법인의 간행물은 일반적인 정보제공 및 지식전달을 위하여 제작된 것으로, 구체적인 회계이슈나 세무이슈 등에 대한 삼일회계법인의 의견이 아님을 유념하여 주시기 바랍니다. 본 간행물의 정보를 이용하여 문제가 발생하는 경우 삼일회계법인은 어떠한 법적 책임도 지지 아니하며, 본 간행물의 정보와 관련하여 의사결정이 필요한 경우에는, 반드시 삼일회계법인 전문가의 자문 또는 조언을 받으시기 바랍니다.

S/N: 2606W-RP-088

©2026 Samil PwC. All rights reserved. PwC refers to the PwC network and/or one or more of its member firms, each of which is a separate legal entity. Please see [pwc.com/structure](https://www.pwc.com/structure) for further details.