

AI 조직 재설계

AI 에이전트 시대, 일하는 방식과 조직 운영을 다시 설계하는 책임입니다.

김경진 변호사

Korean PDF edition

AI 조직 재설계

김경진 변호사

목차에서 서문, 6개 장, 마지막 글로 순서대로 이동할 수 있습니다.

목차

서문. 이 책의 목적과 구성

제1장. 다가오는 AI 에이전트 시대, 일하는 방식이 완전히 달라진다

제2장. AI 도입의 성패를 가르는 10:20:70 법칙과 업무 혁신

제3장. 딱딱한 기계 같은 조직에서, 살아 숨 쉬는 생물 같은 조직으로

제4장. 미래 직장의 새로운 인재와 직업의 재탄생

제5장. 성공을 이끄는 리더의 자세와 조직 문화 바꾸기

제6장. 작은 실험을 회사 전체로 넓히는 전략

마지막 글. 칼을 잡을 것인가, 전쟁을 설계할 것인가

서문. 이 책의 목적과 구성

이 글은 AI 에이전트 시대에 조직이 무엇을 바꾸어야 하는지 설명하기 위해 작성되었다.

AI 도입은 단순한 기술 구매가 아니다. 업무 방식, 의사결정 구조, 인재의 역할, 리더십, 조직문화, 확산 전략을 함께 바꾸는 조직 재설계의 문제다.

이 책은 그 문제를 여섯 개의 흐름으로 정리한다. 제1장은 AI 에이전트가 무엇이며 일하는 방식을 어떻게 바꾸는지 설명한다. 제2장은 AI 도입의 성패를 가르는 10:20:70 법칙과 프로세스 혁신을 다룬다. 제3장은 기계형 조직에서 지능형 조직으로 이동하는 구조적 변화를 설명한다. 제4장은 미래 직장의 인재상과 직업의 변화를 살핀다. 제5장은 리더십과 조직문화의 문제를 다룬다. 제6장은 파일럿 실험을 넘어 전사적 확산으로 가기 위한 전략을 제시한다.

핵심 질문은 하나다. AI가 더 많은 일을 맡는 시대에 사람과 조직은 무엇을 다시 설계해야 하는가.

이 글은 기술 전문가만을 위한 문서가 아니다. 조직을 운영하는 사람, 현장에서 일하는 사람, 앞으로 직업 세계에 들어갈 사람이 AI 변화를 자기 문제로 이해할 수 있도록 구성했다.

2026년 5월.

KIMKJ.COM

제1장. 다가오는 AI 에이전트 시대, 일하는 방식이 완전히 달라진다

1.1 시키는 일만 하던 AI가, 스스로 생각하고 움직이는 AI로 바뀌고 있다

로마 제국이 지중해 전역을 다스리던 시절, 황제에게는 두 부류의 부하가 있었다. 하나는 전령(傳令)이다. 황제가 “갈리아에 편지를 전하라”고 명령하면 그대로 말을 달려 전달하고 돌아왔다. 편지 내용이 맞는지, 갈리아의 상황이 어떤지는 전령의 관심사가 아니었다. 시킨 일만 했다.

다른 하나는 총독이다. 총독은 달랐다. 황제로부터 “속주를 안정시켜라”는 큰 방향만 받으면, 세금을 어떻게 걷을지, 반란의 조짐은 없는지, 도로는 어디에 놓을지를 스스로 판단하고 실행했다. 밤새 보고서를 읽고, 현지 사정을 살피고, 필요하면 군단을 이동시켰다. 황제가 잠든 사이에도 총독의 일은 계속되었다.

지금 우리가 겪고 있는 AI의 변화가 딱 이 모양새다.

지난 20년 동안 우리가 써온 기술은 전령에 가까웠다. 검색 엔진이 그랬다. 네이버나 구글에 “서울 날씨”라고 치면 결과를 보여주는 것까지가 검색 엔진의 몫이었다. 그 결과를 보고 우산을 챙길지 말지는 사람이 직접 판단해야 했다. 정보를 모으고, 비교하고, 결정하고, 실행하는 모든 과정이 사람의 일이었다.

그 뒤에 등장한 챗봇, 그러니까 ChatGPT나 클로드 같은 대화형 AI는 조금 더 똑똑한 전령이었다. “이 보고서를 요약해줘”라고 하면 깔끔하게 요약해주었다. “영어 이메일을 한국어로 번역해줘”라고 하면 번역해주었다. 하지만 여전히, 사람이 먼저 입을 열어야만 움직였다. 명령이 없으면 아무것도 하지 않는, 말하자면 ‘수동적인 도구’였다. 이런 AI를 코파일럿(Copilot)이라고 부른다. 비행기의 보조종사처럼 옆에 앉아서 돕기는 하지만, 기장이 조종간을 잡지 않으면 비행기를 띄울 수 없는 존재다.

그런데 2025년 말부터 완전히 다른 종류의 AI가 나타나기 시작했다.

이 새로운 AI는 전령이 아니라 총독에 가깝다. 사람이 “다음 주 출장 준비 좀 해줘”라고 한마디만 던지면, 이 AI는 혼자서 비행기 시간표를 찾아보고, 호텔을 비교하고, 회의 일정을 조율하고, 필요한 서류를 정리한다. 사람이 하나하나 지시하지 않아도, 스스로 다음에 뭘 해야 할지 판단하고 움직인다. 이런 AI를 ‘에이전트(Agent)’라고 부르고, 이 방식을 ‘에이전틱 AI(Agentic AI)’라고 부른다. 에이전트(Agent)는 영어로 ‘대리인’이라는 뜻이다. 내가 직접 하지 않아도 나를 대신해서 일을 처리해주는 존재. 로마의 총독이 황제를 대신해 속주를 다스렸듯이.

2026년 지금, 이 에이전트 AI는 이미 실험실을 벗어나 현실 세계에서 일하고 있다. 가트너(Gartner)라는 유명한 기술 분석 회사는 2026년 보고서에서 에이전틱 AI가 “기대의 정점”에 올라 있다고 발표했다. 세일즈포스(Salesforce)라는 기업용 소프트웨어 회사는 2026년 4월, 자기 회사의 모든 기능을 AI 에이전트가 직접 조작할 수 있도록 100개가 넘는 새 도구를 공개했다. 사람이 화면을 클릭할 필요 없이, AI가 뒷단에서 알아서 고객 정보를 관리하고 마케팅을 실행하는 구조다. 엔비디아(NVIDIA)도 같은 해 3월에 AI 에이전트를 만들기 위한 도구 모음을 내놓았고, 어도비(Adobe), SAP, 시스코(Cisco) 같은 거대 기업 열다섯 곳이 바로 도입하겠다고 나섰다.

변화의 속도가 무섭다. 미국의 코딩 도구 커서(Cursor)는 2026년 3월, AI 에이전트가 혼자서 백그라운드에서 코드를 점검하고 수정하는 기능을 내놓았다. 개발자가 잠든 사이에도 AI가 일한다. 이 회

사의 연매출은 20억 달러, 우리 돈으로 약 2조 8천억 원을 넘어섰는데, 불과 3개월 만에 매출이 두 배로 뛰었다.

진통제에서 면역 체계로

이 차이를 몸에 비유하면 이해가 쉽다.

코파일럿 형태의 AI는 진통제다. 머리가 아플 때 약을 먹으면 통증이 사라진다. 하지만 약을 먹지 않으면 아무 일도 일어나지 않는다. 내가 "아프다"고 말해야만, 그때서야 작동하는 일회용 해결책이다.

에이전트 AI는 면역 체계에 가깝다. 우리 몸의 면역 세포는 24시간 쉬지 않고 몸 안을 돌아다니며 이상한 바이러스나 세균이 없는지 감시한다. 내가 아프다고 느끼기도 전에 적을 발견하고, 스스로 백혈구를 보내서 싸운다. 잠을 자는 동안에도 일한다. 아무도 시키지 않았는데 알아서 한다.

다른 비유도 가능하다. 기존의 AI는 아주 좋은 '망치'였다. 힘이 세고 정확하지만, 사람이 손에 쥐고 못을 겨냥해서 내리쳐야만 쓸모가 있었다. 에이전트 AI는 '장인(匠人)'이다. "책장 하나 만들어줘"라고 말하면, 장인이 알아서 나무를 고르고, 치수를 재고, 톱질을 하고, 못을 박아서 완성된 책장을 가져다준다.

에이전트를 움직이는 엔진: 관찰, 생각, 행동, 배움

에이전트 AI가 이렇게 자율적으로 움직일 수 있는 비결은, 네 단계가 돌아가는 바퀴에 있다. 관찰(Observe), 생각(Think), 행동(Act), 배움(Learn). 줄여서 OTAL이라고 부른다.

관찰. AI는 주변 상황을 파악한다. 이메일이 도착했는지, 일정이 바뀌었는지, 데이터에 이상한 점은 없는지. 사람이 눈과 귀로 세상을 살피듯, AI는 텍스트, 센서, 인터넷 연결 등을 통해 지금 무슨 일이 벌어지고 있는지 읽어낸다.

생각. 관찰한 것을 바탕으로 "다음에 뭘 해야 하지?"를 궁리한다. 사람이 경험과 논리를 섞어서 판단하듯, AI도 패턴을 읽어내는 직감과 엄격한 논리를 결합해서 계획을 세운다.

행동. 계획을 실행에 옮긴다. 웹사이트에서 항공권을 예매하거나, 회사 시스템에 데이터를 입력하거나, 슬랙(Slack) 메시지를 보내거나. 생각에 그치지 않고 실제로 무언가를 한다.

배움. 가장 중요한 단계다. 행동한 결과를 돌아보고 "잘했나, 못했나"를 스스로 평가한다. 실수했으면 다음번에는 다른 방법을 시도한다. 이것 가능하게 하는 것이 AI의 기억 시스템이다. 지금 하고 있는 작업을 잠시 저장하는 '단기 기억', 과거에 무슨 일을 했고 결과가 어땠는지 일기장처럼 기록하는 '경험 기억', 세상에 대한 전반적인 지식을 담는 '지식 기억'. 이 세 가지 기억이 합쳐져서 에이전트는 일을 할수록 조금씩 더 나아진다.

로마의 총독도 그랬다. 처음 부임했을 때는 속주의 사정을 잘 몰랐지만, 세금을 걷어보고, 반란을 겪어보고, 도로를 놓아보면서 점점 노련해졌다. 경험이 쌓이면 판단이 빨라지고, 판단이 빨라지면 통치가 안정되었다. 에이전트 AI의 학습도 같은 원리다.

두 가지 철학의 충돌: 코파일럿 vs 코워크

이 변화를 바라보는 시각은 크게 둘로 갈린다.

마이크로소프트(Microsoft)는 '코파일럿(Copilot)' 노선을 택했다. 이름 그대로 '부조종사'다. 워드(Word), 엑셀(Excel), 파워포인트(PowerPoint) 같은 기존 프로그램 안에 AI를 심어놓고, 사용자가 일할 때 옆에서 도와주는 방식이다. 안전하고 안정적이다. 하지만 사용자가 한 단계씩 지시를 내려야 한다는 점은 달라지지 않았다.

앤스로픽(Anthropic)은 다른 길을 걸었다. 2025년에 공개한 '클로드 코워크(Claude Cowork)'는 AI를 '새로 채용한 직원'처럼 다룬다. 이 AI 직원은 사용자의 컴퓨터 바탕화면, 파일 폴더에 직접 접근한다. 슬랙(Slack), 세일즈포스(Salesforce), 아사나(Asana) 같은 서로 다른 업무 도구를 넘나들며, 사람이 자리를 비운 사이에도 혼자서 계획을 세우고 복잡한 일을 처리한다.

코파일럿은 "사람이 주인이고, AI는 보조"라는 철학이다.

코워크는 "AI도 한 명의 팀원이다"라는 철학이다.

어느 쪽이 맞는지는 아직 아무도 모른다. 로마가 공화정과 제정 사이에서 오래 흔들렸듯이, 이 문제도 한동안 답이 나지 않을 것이다. 다만 한 가지는 분명하다. 두 방식 모두, AI가 지금보다 훨씬 더 많은 일을 맡게 된다는 방향을 가리키고 있다.

에이전트가 일하려면 필요한 것: MCP와 스킬 파일

에이전트가 회사 안에서 제대로 일하려면, 회사의 여러 시스템에 접속할 수 있어야 한다. 캘린더를 확인하고, 고객 데이터베이스를 열어보고, 회사 메신저에 메시지를 보내야 한다. 그런데 시스템마다 문을 여는 방법이 다르다. 열쇠 모양이 제각각인 셈이다.

여기서 등장하는 것이 MCP(Model Context Protocol)다. 쉽게 말하면 '만능 열쇠'다. 어떤 시스템이든 이 열쇠 하나로 문을 열 수 있게 해주는 표준 규격이다. AI가 회사 방화벽을 뚫지 않고도, 허락받은 범위 안에서 캘린더나 데이터베이스에 안전하게 접근할 수 있다. 2026년 현재, 세일즈포스는 MCP 도구만 60개 넘게 새로 만들었다.

'스킬 파일(Skill.md)'이라는 것도 있다. 이건 에이전트를 위한 업무 지침서, 쉽게 말하면 '사원 수첩'이다. "이런 상황에서는 이렇게 해라", "모르는 건 지어내지 말고 사람에게 물어봐라", "보고서는 이 양식에 맞춰서 써라" 같은 규칙이 적혀 있다. 에이전트는 이 수첩을 읽고 그대로 따른다.

로마 제국이 광대한 영토를 다스릴 수 있었던 비결 중 하나는 '로마법'이라는 통일된 규칙 체계와, 제국 전역을 잇는 도로망이었다. MCP는 AI 세계의 도로, 스킬 파일은 AI 세계의 법전인 셈이다.

사람이 앱을 여는 시대에서, AI가 앱을 여는 시대로

여기서 정말 큰 변화가 일어난다.

지금까지는 사람이 직접 프로그램을 열었다. CRM(고객관리 프로그램)에 로그인해서 고객 정보를 입력하고, ERP(경영관리 프로그램)에 들어가서 재고를 확인하고, 이메일 프로그램을 열어서 거래처에 연락했다. 화면의 버튼을 사람 손가락이 클릭했다.

그런데 에이전트 AI가 등장하면, 사람 대신 AI가 이 프로그램들을 조작한다. 사람은 "이번 달 판매 보고서 만들어줘"라고 말만 하면, AI가 뒷단에서 여러 프로그램을 돌아다니며 데이터를 모으고, 분석하

고, 보고서를 완성한다.

이렇게 되면 화면 디자인, 그러니까 사용자가 보는 화면(UI)의 가치가 떨어진다. 어차피 AI가 쓰는 건데, 예쁜 버튼이 무슨 소용인가. 대신 API, 그러니까 프로그램끼리 데이터를 주고받는 통로의 가치가 올라간다. AI에게는 예쁜 화면이 아니라 잘 정리된 데이터 통로가 필요하기 때문이다.

이 변화의 규모는 상상 이상이다. 2026년 2월, 소프트웨어 회사들의 주가가 하루 만에 약 2,850억 달러(우리 돈 400조 원 가까이)가 증발했다. 서비스나우(ServiceNow)가 7%, 세일즈포스가 7%, 인튜이트(Intuit)가 11% 떨어졌다. AI가 너무 잘 작동해서, 회사들이 더 이상 소프트웨어 구독을 여러 개 살 필요가 없어졌기 때문이다. 가트너는 2030년까지 기업용 소프트웨어의 35%가 AI 에이전트로 대체될 거라고 내다봤다.

호주의 디자인 회사 캔바(Canva)도 주목할 만하다. 2025년에 연매출 40억 달러(약 5조 6천억 원)를 달성한 이 회사는, 자기네 내부에서 쓰던 수많은 업무용 소프트웨어를 버리고 AI 기반 시스템으로 교체했다. 2026년 4월에는 아예 “우리는 더 이상 디자인 도구에 AI를 붙인 회사가 아니라, AI 회사에 디자인 도구를 붙인 것”이라고 선언했다.

기존의 소프트웨어 회사 입장에서 보면, 이걸 시장 자체가 뒤집히는 사건이다. 전에는 직원 수만큼 소프트웨어 이용권을 팔면 됐다. 직원이 100명이면 100개의 자리값을 받았다. 그런데 AI 에이전트가 한 사람 몫의 일을 다섯 명분 해치우면, 회사는 자리값 다섯 개 대신 하나만 사면 된다. 소프트웨어 회사의 수입이 줄어든 수밖에 없다.

로마 제국에서 도로가 깔리자, 도로 위의 무역이 폭발적으로 늘어난 대신 도로 없이 장사하던 지역 상인들은 도태되었다. 기술의 인프라가 바뀌면, 그 위에서 먹고살던 사람들의 운명도 바뀐다. 지금 일어나고 있는 일이 정확히 그것이다.

1.2 일을 직접 하는 사람에서, AI를 지휘하는 사람으로

체스판의 새로운 규칙

체스에서 승패를 가르는 것은 무엇일까.

말(Piece)의 수? 아니다. 양쪽 다 처음에 같은 수의 말을 받는다. 말의 이동 능력? 그것도 아니다. 권의 움직임은 양쪽이 동일하다. 결국 승패를 결정짓는 것은 말을 어디에 놓고, 어떤 순서로 움직이느냐는 전략이다. 말 자체가 아니라, 말을 쥔 사람의 머릿속에 답이 있다.

AI 에이전트 시대에도 같은 일이 벌어진다.

조만간 모든 사람이 비슷하게 뛰어난 AI 에이전트를 갖게 될 것이다. 누구나 강력한 말을 손에 쥔다. 그때 차이를 만드는 것은 AI 자체가 아니라, 그 AI를 어떻게 배치하고 어떤 목표를 향해 움직이게 하느냐 하는 사람의 판단이다.

과거에는 사람이 직접 일을 했다. 보고서를 쓰고, 데이터를 정리하고, 이메일을 보냈다. ‘실행자’가 곧 일잘러였다. 그런데 AI가 이런 실행을 넘겨받기 시작하면, 사람의 역할이 바뀐다. 직접 손을 움직이는 ‘연주자’에서, 여러 연주자를 이끄는 ‘지휘자(오케스트레이터)’로.

오케스트라의 지휘자는 바이올린을 켜지 않는다. 플루트도 불지 않는다. 하지만 모든 악기가 언제 들어오고, 얼마나 세게 연주하고, 어떤 감정을 담아야 하는지를 결정한다. 지휘자 없는 오케스트라는 아무리 뛰어난 연주자가 모여도 소음밖에 만들지 못한다. 이것이 앞으로 사람에게 요구되는 역할이다. AI라는 연주자들을 지휘해서 하나의 결과물을 만들어내는 능력.

AI 팀을 설계하는 사람: 에이전트 매니저

회사에서 일하는 모습을 상상해보자.

옛날 방식이라면, 팀장이 부하 직원 한 명에게 모든 일을 시켰다. "시장 조사도 하고, 보고서도 쓰고, 발표자료도 만들어." 그 직원이 아무리 유능해도, 한 사람이 감당할 수 있는 일에는 한계가 있었다.

에이전트 시대의 팀장은 다르게 일한다. AI 에이전트 여러 개를 각각 다른 역할에 배치한다. 하나는 인터넷을 돌아다니며 최신 뉴스와 시장 정보를 모으는 '정보 수집 에이전트'. 다른 하나는 모아온 정보를 바탕으로 초안을 쓰는 '작성 에이전트'. 또 다른 하나는 회사의 말투와 규칙에 맞게 글을 다듬는 '편집 에이전트'. 팀장은 이 AI 팀이 일을 잘하고 있는지 확인하고, 최종 결과물에 도장을 찍는 역할만 한다.

이걸 정리한 것이 RPM이라는 틀이다.

R은 Role, 역할이다. 이 에이전트는 뭘 하는 녀석인가. 정보를 모으는 녀석인가, 글을 쓰는 녀석인가, 교정을 보는 녀석인가.

P는 Project, 프로젝트다. 이번에 완성해야 할 목표가 뭔가. 예를 들어 "6월 신제품 발표 보도자료 완성"처럼 구체적이어야 한다.

M은 Method, 방법이다. 어떤 규칙을 지켜야 하는가. "고객 이름은 반드시 가명을 써라", "수치는 출처를 밝혀라" 같은 것들이다.

이렇게 하면 AI 에이전트는 사람처럼 역할 분담을 하고, 서로 결과물을 주고받으며 일을 완성한다. 사람은 전체 그림을 그리고, 중간중간 방향을 잡아주면 된다.

반인마(半人馬)의 탄생

그리스 신화에 '켄타우로스(Centaur)'라는 존재가 나온다. 윗몸은 사람이고 아랫몸은 말이다. 사람의 지혜와 말의 힘을 동시에 가진 존재. AI 시대에 가장 값어치 있는 인재가 바로 이 켄타우로스를 닮았다.

프로그래머를 예로 들어보겠다.

예전의 프로그래머는 코드를 처음부터 끝까지 손으로 짰다. 한 줄 한 줄 정성 들여 쓰는 것이 실력의 증거였고 자부심이었다. 장인이 가구를 손으로 깎듯이.

그런데 2026년의 현실은 다르다. 구글의 한 엔지니어는 앤스로픽의 클로드 코드(Claude Code)를 쓰고 나서 "작년에 우리 팀이 1년 걸려 만든 것을 한 시간 만에 뽑아냈다"고 말했다. AI가 반복적이고 기계적인 코드를 순식간에 만들어낸다.

이 상황에서 두 부류의 사람이 갈린다.

한 부류는 AI가 만든 코드의 작은 실수를 비웃으며 “역시 사람 손을 못 따라오지”라고 말하는 사람이다. 이 사람은 자기 손기술에 자부심을 느끼지만, 점점 시대에 뒤쳐진다.

다른 한 부류는 반복적인 코드 작성은 AI에게 맡기고, 자신은 전체 시스템을 어떻게 설계할지, 사용자가 무엇을 원하는지, 이 프로그램이 사업에 어떤 가치를 주는지를 고민하는 데 시간을 쓰는 사람이다. 사람의 머리와 AI의 손을 결합한 켄타우로스.

시장이 원하는 인재는 후자다.

사람이 반드시 끼어들어야 하는 이유

AI가 이렇게 많은 일을 넘겨받으면, 사람은 편해지기만 할까. 그렇지 않다. 오히려 사람의 책임이 더 무거워진다.

AI는 빠르다. 놀라울 만큼 빠르게 데이터를 처리하고 결과를 내놓는다. 하지만 AI는 가끔 거짓말을 한다. 없는 사실을 있는 것처럼 그럴듯하게 지어내는 현상을 ‘환각(Hallucination)’이라고 부른다. 숫자를 엉뚱하게 만들어내거나, 존재하지 않는 법률 조항을 인용하기도 한다.

또 AI에게는 양심이 없다. 이 결정이 누군가에게 피해를 주는지, 이 보고서가 편향되어 있는지를 스스로 판단하지 못한다. 효율의 논리만 있고, 윤리의 감각은 없다.

그래서 ‘사람이 고리 안에 있어야 한다(Human-in-the-loop)’는 원칙이 중요해진다. AI가 내놓은 결과를 사람이 검토하고, 문제가 있으면 되돌리고, 최종 승인을 내리는 과정. 이것을 빼면, 빠르지만 위험한 기계가 회사를 좌우하게 된다.

로마사에서도 비슷한 교훈을 찾을 수 있다. 로마의 군단(Legion)은 당대 최강의 전투 기계였다. 하지만 군단을 어디로 보낼지, 언제 싸우고 언제 멈출지를 결정하는 것은 군단장이나 원로원이었다. 군단 자체가 전쟁을 결정한 적은 없다. 무기가 아무리 강해도, 그것을 쥔 손과 머리가 없으면 재앙이 된다.

AI도 마찬가지다. AI가 해방시켜주는 것은 반복적인 작업에 쓰이던 시간이다. 대신 사람에게 돌아오는 것은 “결과에 대한 책임”이라는, 어쩌면 더 무거운 짐이다.

새로 생겨나는 직업들

이 거대한 변화는 새로운 종류의 일자리를 만들어내고 있다.

‘AI 제품 매니저(AI PM)’라는 직무가 대표적이다. 이 사람은 통역사 같은 역할을 한다. 회사의 사업 부서가 “고객 만족도를 올리고 싶어”라고 말하면, 그 모호한 바람을 AI가 알아들을 수 있는 구체적인 업무 지시와 작업 흐름으로 번역해준다. 사람의 말을 기계의 말로 바꾸고, 기계의 결과물을 사람이 이해할 수 있는 형태로 다시 바꾸는 일이다.

‘AI 윤리 엔지니어’도 생겨나고 있다. AI가 편향된 데이터를 학습해서 특정 집단에 불리한 결과를 내놓지는 않는지, AI의 판단 과정이 투명한지를 감시하는 사람이다. 그리고 AI가 만들어낸 결과물의 품질을 평가하는 ‘AI 산출물 평가관’이라는 직무도 등장하고 있다.

공통점이 있다. 이 새로운 직업들은 모두, AI를 잘 쓰는 능력과 사람만이 가진 판단력을 함께 요구한다는 것이다.

미래의 직장에서 가장 값진 사람은, 엑셀을 빨리 다루는 사람도 아니고 코딩을 잘하는 사람도 아닐 것이다. AI가 무엇을 할 수 있고 무엇을 할 수 없는지를 정확히 알면서, 기계가 못하는 것(공감, 직관, 윤리적 판단, 창의적 발상)을 업무 과정에 녹여낼 줄 아는 사람. 그런 사람이 가장 대접받는 시대가 오고 있다.

카이사르는 갈리아 전쟁에서 직접 칼을 휘두른 적이 거의 없다고 전해진다. 그의 천재성은 누구를 어디에 배치하고, 어떤 순서로 공격하며, 언제 멈출지를 아는 데 있었다. 2026년, 우리에게도 비슷한 질문이 던져졌다. 칼을 휘두르는 법을 배울 것인가, 아니면 전쟁을 설계하는 법을 배울 것인가.

KIMKJ.COM

제2장. AI 도입의 성패를 가르는 10:20:70 법칙과 업무 혁신

2.1 성공의 비밀은 기술이 아니라 사람에게 있다

90%가 뛰어들었는데, 9%만 살아남았다

로마가 지중해의 패권을 잡기까지, 수백 개의 도시 국가가 같은 바다에서 경쟁했다. 군선을 가진 나라는 많았다. 갑옷을 만드는 기술도 비슷비슷했다. 그런데 결국 지중해를 손에 넣은 것은 로마 하나였다. 배와 갑옷의 차이가 아니었다. 그것들을 쓰는 사람과 조직의 차이였다.

지금 AI 세계에서 똑같은 일이 벌어지고 있다.

전 세계 기업의 90%가 이미 AI를 도입했다. 회의실에서 “우리도 AI 씁니다”라고 말 못 하는 회사가 거의 없는 셈이다. 그런데 AI를 써서 실제로 돈을 벌거나 일하는 방식을 바꾸는 데 성공한 회사는 겨우 9%다. 열 곳 중 아홉 곳이 경기장에 뛰어들었는데, 결승선을 통과한 건 열 곳 중 한 곳도 안 된다.

왜 이런 일이 생길까.

대부분의 회사가 같은 실수를 저지르기 때문이다. “비싼 AI를 사면 된다”는 착각. 시장에서 가장 좋다는 AI 프로그램을 거금을 들여 구매하면, 마치 마법처럼 회사가 변할 거라고 믿는 것이다. 이건 마치 세계에서 제일 비싼 F1 경주차를 샀다고 해서, 면허도 없는 사람이 갑자기 프로 레이서가 될 거라고 기대하는 것과 같다. 차가 아무리 좋아도, 운전하는 사람이 준비되어 있지 않으면 출발선에서 벽을 들이받는다.

10:20:70이라는 숫자의 의미

보스턴 컨설팅 그룹(BCG)이라는, 세계에서 가장 유명한 경영 자문 회사가 있다. 이 회사가 수많은 기업의 AI 도입 사례를 분석하고 나서 하나의 법칙을 내놓았다. ‘10:20:70 법칙’이다.

AI 도입이 성공할 때, 그 성공의 10%는 AI 프로그램 자체의 성능에서 나온다.

20%는 컴퓨터, 서버, 데이터 저장소 같은 기술 장비에서 나온다.

나머지 70%는 사람과 일하는 방식의 변화에서 나온다.

10%와 20%를 합치면 30%다. 이 30%는 뭘까. 돈만 있으면 살 수 있는 것들이다. 좋은 AI 모델? 구매할 수 있다. 빠른 컴퓨터? 주문하면 된다. 클라우드 서버? 카드 결제하면 당장 쓸 수 있다.

2026년 현재, 뛰어난 AI 모델은 오픈AI(OpenAI), 앤스로픽(Anthropic), 구글(Google) 등 여러 회사가 내놓고 있어서, 누구나 비슷한 수준의 AI를 쓸 수 있게 되었다. MIT 슬론 경영대학원의 연구도 같은 결론을 내렸다. AI 모델은 빠르게 ‘누구나 쓸 수 있는 공공재’가 되어가고 있다고.

그렇다면 진짜 승부는 어디서 갈리느냐. 나머지 70%다. 사람들이 AI를 어떻게 받아들이느냐, 일하는 순서를 어떻게 바꾸느냐, 조직 문화가 새로운 도구를 허용하느냐. 이것들은 돈으로 살 수 없다. 시간과 노력과 리더의 결단이 필요하다.

로마가 강했던 이유를 떠올려보면 이해가 된다. 로마의 검(글라디우스)은 특별히 뛰어난 무기가 아니었다. 짧고 투박했다. 켈트족의 긴 칼이 겉보기에는 더 위협적이었다. 하지만 로마 군단이 이겼다. 검이 좋아서가 아니라, 그 검을 쓰는 군단 편제, 훈련 체계, 지휘 구조가 뛰어났기 때문이다. 무기의 성능이 30%라면, 그것을 쓰는 조직과 사람이 70%였다. 2천 년 전이나 지금이나, 이 비율은 놀라울 만큼 비슷하다.

회사가 못 따라가면, 직원이 혼자 달린다

마이크로소프트와 링크드인이 전 세계 직장인 2만 명 이상을 조사한 보고서가 있다. 그 결과가 흥미롭다. AI를 잘 쓰는 사람과 못 쓰는 사람의 차이에서, 개인의 능력이 영향을 미치는 비율은 32%에 그쳤다. 나머지 67%는 회사의 문화와 지원 체계가 결정했다.

아무리 똑똑한 사람이 혼자 AI를 열심히 배워도, 회사가 “그건 아직 쓰지 마” “보안 문제가 있어”라며 막아버리면 소용이 없다는 뜻이다.

그런데 재미있는 현상이 벌어지고 있다. 회사가 막든 말든, 직원들은 이미 AI를 쓰고 있다. 78%의 직장인이 회사의 허락 없이 개인적으로 챗GPT나 클로드 같은 AI 도구를 업무에 쓰고 있다는 조사 결과가 나왔다. 이것을 ‘새도우 AI(Shadow AI)’라고 부른다. 그림자처럼 회사 모르게 쓴다는 뜻이다.

2026년 기준으로 상황은 더 심각해졌다. 98%의 기업에서 허가 없이 AI를 쓰는 직원이 발견되고 있다. 직원의 47%는 개인 계정으로 AI 도구에 접속해서 회사 일을 처리한다. 38%는 회사의 기밀 정보를 AI에 넣어본 적이 있다고 답했다. IBM의 2025년 보고서에 따르면, 이런 새도우 AI 때문에 정보가 유출되면 한 건당 평균 67만 달러(약 9억 4천만 원)의 추가 비용이 든다.

직원들이 규칙을 어기려고 이러는 걸까. 아니다. 일을 빨리, 잘하고 싶어서 그러는 거다. 회사가 좋은 AI 도구를 제공하지 않으니, 직원이 알아서 찾아 쓰는 것이다. 실제로 한 의료 기관에서 공식 AI 도구를 제공했더니, 허가 없이 AI를 쓰는 비율이 89%나 줄어들었다는 결과가 있다. 막는 게 답이 아니라, 제대로 된 도구를 내주는 게 답이라는 뜻이다.

가장 안타까운 부류가 있다. 개인적으로는 AI를 아주 잘 쓰는데, 회사의 지원이 없어서 성장이 막힌 사람들이다. 이 사람들은 주말에 시간을 쪼개서 AI 사용법을 공부한다. 일을 열 배 빠르게 처리하는 방법을 알아낸다. 하지만 그걸 동료들과 나누지 않는다. 왜? 회사에서 알아주지도 않고, 오히려 “쓸데없는 짓 한다”며 눈총을 받을까 봐. 혁신을 시도했다가 욕먹는 게 두렵기 때문이다.

이 현상은 2천 년 전 로마에서도 있었다. 유능한 장군이 전장에서 새로운 전술을 시도하고 싶어도, 원로원이 “전통대로 하라”고 고집하면 꼼짝할 수 없었다. 한니발이 알프스를 넘어 이탈리아를 침공했을 때, 로마 장군 파비우스 막시무스(Fabius Maximus)는 정면 대결을 피하고 지연 전술을 택했다. 원로원과 시민들은 그를 ‘꾸물이(Cunctator)’라고 비웃었다. 새로운 방법을 시도하는 사람은 언제나 기존 체계의 저항에 부딪힌다.

AI 시대에도 마찬가지다. 기술의 문제가 아니다. 사람과 문화의 문제다.

실패를 허락하는 조직이 이긴다

AI를 도입할 때 기술 부서에만 맡겨놓고, 직원들의 마음은 신경 쓰지 않는 회사가 많다. 이러면 ‘문화 빚(Culture Debt)’이 쌓인다. 빚은 갚지 않으면 이자가 붙듯이, 문화 빚도 방치하면 점점 커진다.

56%의 기업이 AI를 도입하면서 직원들의 불안감이나 공정성 같은 문제를 무시했다. 그 결과, 60% 이상의 직원이 '변화 피로(Change Fatigue)'를 느끼고 있다. 새 도구가 나올 때마다, 새 시스템이 바뀔 때마다, 지쳐서 더 이상 따라가기 싫어지는 것이다.

이 문제의 열쇠가 '심리적 안전감(Psychological Safety)'이다. 어려운 말처럼 들리지만, 뜻은 간단하다. "실수해도 괜찮은 분위기." 새로운 AI 도구를 써보다가 잘못해도 혼나지 않고, 바보 취급당하지 않는 환경. "한번 해봤는데 안 됐어요"라고 솔직하게 말할 수 있는 직장.

이건 단순히 '직원 복지' 차원의 이야기가 아니다. 돈과 직결되는 문제다. 직원들이 실수를 두려워하지 않고 새로운 시도를 할 수 있는 팀에서는, 새로운 아이디어가 27% 더 많이 나오고, 결정을 내리는 속도가 35% 빨라지며, 회사를 그만두는 사람이 50%나 줄어든다는 연구가 있다.

스페이스X(SpaceX)를 예로 들어보겠다. 일론 머스크(Elon Musk)의 이 우주 회사는 로켓이 폭발해도 "실패"라고 부르지 않는다. "빠른 비계획 해체(Rapid Unscheduled Disassembly)"라는 표현을 쓴다. 농담 같지만, 이 말에는 철학이 담겨 있다. 실패는 부끄러운 것이 아니라, 다음번에 더 잘하기 위한 데이터다. 로켓이 터지면 "왜 터졌지?"를 분석하고, 그 교훈으로 다음 로켓을 더 낫게 만든다.

BCG의 에르네스토 파가노(Ernesto Pagano) 수석 파트너도 같은 말을 했다. "직원들이 AI를 자기를 대체하는 도구가 아니라, 자기를 돕는 도구로 이해하게 되면 비로소 받아들입니다." 강요해서 되는 일이 아니다. 신뢰를 쌓아야 한다.

기술자와 사업가 사이의 통역사: AI 제품 매니저

70%의 변화를 이끌려면 새로운 종류의 인재가 필요하다.

지금까지 회사에는 큰 골짜기가 하나 있었다. 한쪽에는 기술 부서가 있다. 컴퓨터와 AI를 다루는 사람들이다. 다른 한쪽에는 사업 부서가 있다. 고객을 만나고, 물건을 팔고, 돈을 버는 사람들이다. 이 두 부서는 같은 회사에서 일하면서도 서로 말이 잘 통하지 않았다.

사업 부서가 "고객 만족도를 올리고 싶어"라고 말하면, 기술 부서는 "그게 구체적으로 무슨 뜻이죠? 어떤 수치를 올리라는 건가요?"라고 되묻는다. 사업 부서는 기술의 말을 모르고, 기술 부서는 사업의 말을 모른다. 이 소통 실패 때문에 엉뚱한 AI가 만들어지고, 돈과 시간이 낭비된다.

여기서 등장하는 것이 'AI 제품 매니저(AI PM)'다. 이 사람은 통역사다. 사업 부서의 모호한 바람("고객이 더 만족하게 해주세요")을 기술 부서가 알아들을 수 있는 구체적인 업무 지시("고객 문의 응답 시간을 현재 24시간에서 2시간으로 줄이는 AI 챗봇을 만드세요. 단, 오답률은 5% 이하로 유지하세요")로 번역해주는 사람이다.

이 사람은 코딩을 직접 하지 않아도 된다. 대신 두 가지를 동시에 이해해야 한다. 사업이 어떻게 돌아가는지, 그리고 AI가 무엇을 할 수 있고 무엇을 못하는지. 두 세계의 말을 모두 구사하는 '이중 언어 구사자'인 셈이다.

로마 제국이 그리스, 이집트, 갈리아 등 수십 개의 다른 문화를 하나로 통합할 수 있었던 이유 중 하나는, 라틴어와 그리스어를 모두 구사하면서 각 지역의 사정을 이해하는 '통역 관료'들이 있었기 때문이다. 기술과 사업 사이에도 이런 통역자가 필요하다. 2026년의 기업에서 가장 몸값이 비싼 인재가 바로 이 AI PM이다.

사장이 먼저 써야 한다

마지막으로, 70%의 변화를 끌어내는 가장 강력한 힘은 무엇일까.

마이크로소프트 보고서의 답은 의외로 간단했다. “리더가 직접 AI를 쓰는 모습을 보여주는 것.” 사장이 회의 시간에 직접 AI에게 질문하고, 결과를 함께 검토하는 모습을 보여주면, 직원들은 “아, 이건 진짜구나”라고 느낀다. 반대로 사장이 뒤에서 보고서만 받으면서 “AI 잘 도입해”라고만 말하면, 직원들은 “그냥 유행이구나, 호지부지되겠지”라고 생각한다.

BCG의 파가노 파트너도 같은 관찰을 전했다. 200쪽짜리 계획서를 만들어놓고 컨설턴트가 떠나면 아무도 안 본다. 대신, 리더가 직접 AI를 쓰면서 팀원들과 함께 시행착오를 겪는 것이 훨씬 효과적이다. 그는 이런 리더를 ‘직접 뛰는 최고경영자(ICCEO, Individual Contributor CEO)’라고 불렀다.

결국, AI 도입의 진짜 승부처는 “어떤 AI를 샀느냐”가 아니다. “우리 회사의 사장과 직원이, AI와 함께 일하는 문화를 얼마나 빨리 만들어냈느냐”에 달려 있다.

2.2 일하는 순서를 처음부터 다시 그려야 한다

52개 프로젝트를 했는데, 달라진 게 없다

미국에 70억 달러(약 10조 원) 규모의 대형 보험 회사가 있었다. 이 회사는 AI에 진심이었다. 무려 52개의 AI 시범 프로젝트를 돌렸다. AI 모델의 정확도도 목표치를 달성했다. 기술적으로는 성공이었다.

그런데 18개월이 지나고 보니, 회사의 일하는 방식은 전혀 바뀌지 않았다.

보험금 심사를 하는 직원은 여전히 예전 방식대로 서류를 넘기고 있었다. AI는 한쪽 구석에서 가끔 참고하는 별도의 화면으로 방치되어 있었다. 52개 프로젝트에 들인 돈과 시간은 거의 허공에 뿌린 셈이었다.

왜 이런 일이 벌어졌을까.

기술은 바꿨는데, 일하는 순서는 바꾸지 않았기 때문이다.

이걸 전문가들은 ‘볼트온(Bolt-on)’ 방식이라고 부른다. 볼트(Bolt)는 나사다. 기존의 낡은 기계에 새 부품을 나사로 억지로 붙여놓는 방식. 1980년대에 만든 낡은 자동차 차체에 2026년형 최신 엔진을 박아 넣는 것과 같다. 엔진이 아무리 좋아도, 차체가 그 힘을 감당하지 못하면 달리다가 부서진다.

많은 회사가 이 실수를 저지른다. 30년 동안 써온 업무 절차는 그대로 두고, 그 위에 AI만 얹는다. 처음에는 조금 빨라지는 것 같다. 하지만 그 정도의 개선은 경쟁사도 똑같이 할 수 있다. 진짜 승부는 일하는 순서 자체를 처음부터 다시 그리는 회사와, 낡은 순서에 새 도구만 끼워 넣는 회사 사이에서 갈린다.

BCG가 2026년 3월에 발표한 보고서는 이 점을 분명히 했다. AI 에이전트를 대규모로 쓰려면, 기존의 업무 과정 자체를 다시 설계해야 한다고. 잘 설계된 에이전트 기반 업무 과정은 “내가 제대로 들었나?” “빠뜨린 게 없나?” 같은 확인 단계를 크게 줄여준다. 기존에 며칠 걸리던 검토가 몇 분 만에 끝

나기도 한다.

물어야 할 질문이 바뀌어야 한다. "AI를 써서 기존 일을 어떻게 더 빨리 할까?"가 아니라, "AI가 있는 세상에서, 이 일을 처음부터 만든다면 어떻게 만들까?"

매번 사람이 확인해야 하면 AI를 쓰는 의미가 없다

기존 방식에 AI를 억지로 끼워 넣으면 생기는 문제가 하나 더 있다. 결재 병목이다.

많은 회사에서 AI가 보고서를 만들면, 먼저 담당자가 확인한다. 그 다음 팀장이 확인한다. 그 다음 부장이 확인한다. AI가 아무리 빨리 결과를 내놓아도, 사람이 한 단계씩 승인하는 구조가 그대로 남아 있으면 전체 속도는 사람 속도에 묶인다. 이름만 AI이지, 실제로는 예전의 수작업 과정을 그대로 반복하는 꼴이다.

그래서 진짜 변화는 승인 구조 자체를 바꾸는 것이다. 매번 사람이 확인하는 '절차적 안전장치' 대신에, 시스템 설계 단계에서부터 안전장치를 심어놓는 '구조적 안전장치'가 필요하다.

예를 들어보겠다. AI가 회사의 고객 데이터에 접근할 때, "이 AI는 A등급 데이터까지만 볼 수 있고, B등급은 못 본다"는 규칙을 처음부터 시스템에 박아놓는 것이다. 매번 사람이 "이거 봐도 돼?" 하고 확인해줄 필요가 없다. 규칙이 이미 설계에 녹아 있으니까.

목표는 분명하다. 기존에 10단계였던 승인 과정을 AI로 10배 빠르게 만드는 것이 아니다. 10단계의 과정 자체를 4단계로 줄이는 것이다. 과정의 수를 줄이되, 안전은 구조로 보장하는 것.

로마의 도로를 생각하면 이해가 쉽다. 로마가 도로를 깔기 전에는, 물자를 나르는 마차가 마을마다 멈춰서 통행세를 내야 했다. 마을이 열 개면 열 번 멈춰야 했다. 로마는 이 구조를 바꿨다. 주요 도로에는 검문소를 없애고, 대신 도로의 폭과 규격을 표준화해서 어떤 마차든 막힘 없이 달릴 수 있게 했다. 검문소를 빠르게 통과시키는 게 아니라, 검문소 자체를 없앤 것이다. AI 시대의 업무 재설계도 같은 원리다.

1940년대 영국 탄광에서 배우는 교훈

업무를 다시 설계한다고 하면, 기술 이야기만 할 것 같지만 그렇지 않다. 놀랍게도, 이 문제의 답은 80년 전 영국의 탄광에서 나왔다.

1940년대 영국. 탄광에 아주 좋은 새 채굴 기계가 도입되었다. 기계의 성능만 보면, 석탄 생산량이 크게 늘어야 했다. 그런데 실제로는 생산량이 오히려 떨어졌다. 기계가 나쁜 게 아니었다. 문제는 다른 데 있었다. 새 기계를 도입하면서 탄광 노동자들의 팀 구조가 바뀌었는데, 기존에 서로 도우며 일하던 끈끈한 팀워크가 깨져버린 것이다. 기술은 좋아졌는데, 사람들 사이의 관계가 망가지니까 전체 결과가 나빠졌다.

이 사건을 연구한 학자들이 내린 결론이 있다. "기술 시스템과 사람(사회) 시스템을 동시에 바꿔야 한다. 하나만 바꾸면 오히려 망한다." 이걸 '결합 최적화(Joint Optimization)'라고 부른다.

2026년의 AI 도입에도 같은 원리가 적용된다. AI라는 새 기술을 넣으면서, 사람들이 일하는 방식, 팀의 구조, 권한의 분배를 동시에 재설계해야 한다. AI만 넣고 나머지는 그대로 두면, 80년 전 영국 탄

광과 같은 결과가 나온다.

무엇을 AI에게 맡기고, 무엇을 사람이 할 것인가

일을 다시 설계할 때, 네 가지 칸으로 나눠서 생각하면 도움이 된다.

첫 번째 칸은 '말기기'다. 데이터를 입력하거나, 긴 문서를 요약하거나, 매일 같은 형식의 보고서를 만드는 일. 사람이 해도 되지만, AI가 훨씬 빠르고 실수도 적은 일이다. 이런 건 AI에게 통째로 넘긴다.

두 번째 칸은 '돕기'다. AI가 초안을 쓰고, 사람이 다듬는 방식이다. 예를 들어 AI가 고객에게 보낼 이메일 초안을 만들면, 사람이 읽어보고 어색한 부분을 고친다. AI의 속도와 사람의 감각이 합쳐진다.

세 번째 칸은 '증폭'이다. AI가 사람의 능력을 몇 배로 키워주는 경우다. 예를 들어 의사가 환자의 검사 결과를 분석할 때, AI가 수천 건의 비슷한 사례를 순식간에 찾아서 보여준다. 의사의 경험에 AI의 데이터 처리 능력이 더해져서, 혼자서는 생각해내지 못했을 판단을 내릴 수 있게 된다.

네 번째 칸은 '사람만의 영역'이다. 고객의 슬픔에 공감하는 것, 윤리적으로 옳은지 판단하는 것, 최종 책임을 지는 것. 이런 일은 AI가 할 수 없고, 해서도 안 된다. 반드시 사람의 몫으로 남겨둬야 한다.

모든 업무를 이 네 칸에 하나씩 넣어보는 것. 이것이 업무 재설계의 첫걸음이다.

영업 사원의 하루가 이렇게 바뀐다

구체적인 예를 들어보겠다.

회사에서 물건을 파는 영업 사원이 있다. 이 사람의 하루를 조사해보니, 실제로 고객을 만나서 대화하는 시간은 전체의 25%밖에 안 됐다. 나머지 75%는 뭘 했느냐. 고객관리 프로그램(CRM)에 데이터를 입력하고, 제안서 초안을 만들고, 회의 일정을 잡고, 보고서를 쓰는 데 시간을 쏟았다. 정작 돈을 벌어드 주는 일(고객 만남)에 쓰는 시간이 전체의 4분의 1밖에 안 된 것이다.

업무를 재설계하면 이렇게 바뀐다. CRM 데이터 입력? AI 에이전트가 회의 녹음을 자동으로 분석해서 입력한다. 제안서 초안? AI가 고객의 과거 거래 기록과 업종 정보를 분석해서 맞춤형 초안을 만들어 놓는다. 회의 일정? AI가 양쪽 캘린더를 확인해서 가능한 시간을 잡아준다.

그러면 영업 사원이 고객을 만나는 시간이 두 배로 늘어난다. 그리고 이 사람의 역할도 바뀐다. 예전에는 데이터 입력까지 해야 하는 '잡일꾼'이었다면, 이제는 AI가 만든 전략을 검토하고 최종 결정을 내리는 '지휘자'가 된다.

낡은 기계에서 살아있는 생물로

이 모든 변화가 합쳐지면, 회사 자체의 모습이 달라진다.

지금까지의 회사는 '기계'와 비슷했다. 부서마다 벽이 높고, 위에서 아래로 명령이 내려오고, 정해진 절차대로만 움직였다. 톱니바퀴처럼 규칙적이지만, 톱니바퀴는 갑자기 방향을 바꿀 수 없다.

AI 에이전트가 사람과 섞여서 일하는 미래의 회사는 '살아있는 생물'에 가까워진다. 환경이 바뀌면 실시간으로 반응하고, 필요에 따라 형태를 바꾼다. 문어가 좁은 틈을 통과할 때 몸 모양을 자유자재

로 바꾸듯이.

이런 회사가 되려면, 질문 하나를 바꿔야 한다.

“우리가 가진 업무를 AI로 어떻게 개선할까?”

이 질문을 버리고, 이렇게 물어야 한다.

“만약 처음부터 AI가 있는 상태에서 회사를 만들었다면, 이 일을 어떻게 설계했을까?”

예를 들어보겠다. 은행의 대출 심사가 원래 3일 동안 6번의 결재를 거쳐야 했다고 하자. 기존 방식에 AI를 끼워 넣으면, 3일이 2일로 줄어드는 정도다. 하지만 처음부터 다시 설계하면? AI가 서류를 분류하고, 위험 요소를 찾아내고, 판단 근거를 정리한 메모 초안까지 만들어놓는다. 사람은 그 메모를 읽고 최종 승인만 내린다. 3일이 4시간으로 줄어든다.

이 차이가 곧 경쟁력의 차이다. 한쪽은 기존 도로에 아스팔트를 다시 깐 것이고, 다른 한쪽은 도로 자체를 새로 설계한 것이다.

로마가 지중해 세계의 주인이 된 것은, 다른 나라보다 더 좋은 배를 만들어서가 아니었다. 도로를 새로 깔고, 법을 새로 만들고, 군대를 새로 편제하고, 통치 구조를 새로 설계했기 때문이다. 기존의 것을 개선한 게 아니라, 처음부터 다시 만들었다. 2026년, AI 시대에 살아남을 기업도 같은 길을 걸어야 한다.

KIMKJ.COM

제3장. 딱딱한 기계 같은 조직에서, 살아 숨 쉬는 생물 같은 조직으로

3.1 피라미드가 무너지고 있다

100년 된 설계도의 한계

로마 군단은 피라미드였다.

맨 위에 총사령관이 있고, 그 아래 군단장이 있고, 그 아래 대대장이 있고, 그 아래 백인대장이 있고, 맨 아래에 병사들이 있었다. 위에서 명령이 내려오면 아래로 차례대로 전달되었고, 아래에서 올라온 정보는 층층이 걸려져서 위로 올라갔다. 이 구조는 기원전 3세기부터 서기 5세기까지, 700년 넘게 작동했다.

지난 100년 동안의 회사도 같은 모양이었다. 맨 위에 사장, 그 아래 부사장, 그 아래 부장, 과장, 대리, 사원. 피라미드형 조직이다. 일을 가장 작은 단위로 쪼개서 각 부서에 맡기고, 중간 관리자들이 아래의 보고를 받아서 위로 올리고, 위의 지시를 받아서 아래로 내렸다. 세상이 천천히 변하던 시절에는 이 구조가 잘 맞았다. 안정적이고, 예측 가능했다.

그런데 2026년, 이 100년 된 피라미드가 흔들리고 있다.

PwC(프라이스워터하우스쿠퍼스)라는 세계적인 회계 법인의 조사에 따르면, AI를 도입한 기업의 56%가 아무런 효과를 보지 못했다. 더 놀라운 수치가 있다. AI 시범 프로젝트의 95%가 실패했다는 보고도 있다. 왜 그럴까. 좋은 AI를 낳은 피라미드 조직에 억지로 끼워 넣었기 때문이다. 한 연구에 따르면, 첨단 AI를 딱딱한 하향식 조직에 넣으면 생산성이 오히려 30%나 떨어진다.

로마 군단의 피라미드가 700년간 버텼던 이유는, 전쟁의 속도가 느렸기 때문이다. 적이 쳐들어오는데 몇 주, 몇 달이 걸렸다. 보고가 올라오고 명령이 내려오는 데 며칠이 걸려도 괜찮았다. 하지만 만약 적이 빛의 속도로 공격해온다면? 보고서가 대대장, 군단장을 거쳐 총사령관에게 올라가는 동안 전쟁은 이미 끝나 있을 것이다.

AI 시대의 비즈니스가 바로 그런 전장이다. 경쟁자의 공격이 며칠이 아니라 몇 시간 만에 날아온다. 피라미드 꼭대기까지 보고가 올라가고 결정이 내려오기를 기다리는 사이에, 빠른 회사는 이미 시장을 먹어치운다.

정보의 흐름이 완전히 뒤집혔다

마이크로소프트의 CEO 사티아 나델라(Satya Nadella)는 이 변화를 가장 정확하게 집어냈다. 2026년 1월, 스위스 다보스에서 열린 세계경제포럼에서 나델라는 이렇게 말했다. AI가 정보의 흐름 전체를 평평하게 만들어버렸다고. 구조적으로 다시 설계해야 한다고.

과거에는 이런 식이었다. 회사의 각 부서가 자기 분야의 정보를 모았다. 영업팀은 영업 정보, 재무팀은 재무 정보, 고객 지원팀은 고객 정보. 이 정보들은 각 부서의 실무자가 요약해서 팀장에게 올리고, 팀장이 다시 정리해서 부장에게 올리고, 부장이 또 정리해서 임원에게 올렸다. 이 과정에서 몇 주가 걸렸다. 위로 올라가는 동안 정보가 걸려지고, 왜곡되기도 했다. 중간 관리자의 역할은 바로 이것이었다. 아래의 정보를 모아서 위로 올리는 일.

그런데 AI 에이전트가 등장하자, 이 구조가 무너졌다.

나델라 자신이 50개 회의를 준비할 때의 이야기를 했다. 예전이라면 각 부서가 몇 주에 걸쳐 자료를 모으고, 보고서를 작성해서 올려야 했다. 하지만 지금은 AI 에이전트에게 "이 고객에 대해 전체적으로 브리핑해줘"라고 말하면, AI가 영업 데이터, 투자 기록, 고객 지원 내역을 동시에 가로질러서 몇 초 만에 완벽한 요약물을 만들어낸다.

정보가 아래에서 위로 천천히 올라가는 게 아니라, 모든 곳에서 동시에, 순식간에 모인다.

이렇게 되면 중간에서 정보를 모으고 전달하던 사람들의 역할이 사라진다. 보고서를 요약해서 위로 올리고, 위의 결정을 아래로 전달하고, 승인 도장을 찍는 일. AI가 이 모든 것을 대신하기 때문이다.

나델라는 실제로 이 철학을 자기 회사에 적용했다. 2026년 초, 마이크로소프트는 코파일럿(Copilot) 조직을 대대적으로 통합했다. 소비자용과 기업용으로 나뉘어 있던 팀을 하나로 합치고, 제이콥 안드레우(Jacob Andreou)를 부사장으로 올려서 나델라에게 직접 보고하게 만들었다. 부서 사이의 벽을 없앤 것이다.

더 파격적인 조치도 있었다. 나델라는 주간 'AI 가속 회의'를 만들었는데, 이 회의에서 고위 임원은 발표를 하지 못하게 했다. 대신 현장의 엔지니어들이 직접 나서서, 걸러지지 않은 날것의 정보를 CEO에게 전달하도록 바꿨다. 중간 관리자를 건너뛰고, 일하는 사람의 목소리가 바로 꼭대기에 닿도록 한 것이다.

아마존이 피라미드를 부순 방법

아마존의 CEO 앤디 재시(Andy Jassy)는 더 직접적으로 움직였다.

2024년 9월, 재시는 회사 전체에 지시를 내렸다. "각 부서에서 관리자 대비 실무자의 비율을 15% 이상 높여라." 2025년 1분기가 끝나기도 전에 아마존은 이 목표를 이미 달성했다.

재시의 말이 인상적이다. "사람이 많아지면 중간 관리자가 늘어난다. 중간 관리자가 늘어나면, 다들 선의로 모든 일에 자기 도장을 찍으려 한다. 그러면 결정을 내리기 위한 사전 회의, 그 사전 회의를 위한 사전 회의, 그 사전 회의를 위한 사전 회의가 생긴다." 회의를 위한 회의를 위한 회의. 이게 피라미드 조직의 병이다.

재시는 '관료주의 제보 메일함'도 만들었다. 직원 누구나 "이건 쓸데없는 절차입니다"라고 제보할 수 있는 창구다. 1,000통 가까운 제보가 들어왔고, 아마존은 이를 바탕으로 375개 이상의 절차를 바꿨다.

재시의 목표는 분명했다. "세계에서 가장 큰 스타트업처럼 움직이자." 거대한 회사이지만, 빠르고 가벼운 작은 회사처럼 결정을 내리자.

이건 단순히 돈을 아끼려는 구조조정이 아니다. 경쟁사가 이용할 수 있는 '결정의 지연'이라는 약점을 조직도에서 물리적으로 없애버리는 행위다. 전쟁에서 명령 체계가 복잡한 군대는 빠른 적에게 진다. 비즈니스도 다르지 않다.

부서의 벽이 무너진다

피라미드가 위아래로 짓눌리고 있다면, 부서 사이의 벽은 좌우로 허물어지고 있다.

옛날 회사를 상상해보자. 영업팀, 재무팀, 마케팅팀이 각각 자기 방에 앉아 있다. 자기 팀만의 프로그램을 쓰고, 자기 팀만의 목표를 추적한다. 서로 필요한 정보가 있으면, 이메일을 보내거나 회의를 잡아야 했다. 정보가 한 방에서 다른 방으로 넘어가려면 문을 두드리고 허락을 받아야 했다. 이 벽을 '사일로(Silo)'라고 부른다. 곡식을 저장하는 원통형 창고처럼 각 부서가 따로 떨어져 있다는 뜻이다.

AI 에이전트는 이 벽을 뚫고 다닌다.

예를 들어, 광고 회사에서 고객이 새 광고를 요청했다고 하자. 예전이라면 영업 사원이 고객의 요구를 받아서 기획 부서에 넘기고, 기획 부서가 정리해서 제작 부서에 넘기고, 제작 부서가 만든 결과를 다시 영업 사원이 가져와서 고객에게 전달했다. 정보가 부서의 문을 세 번 두드려야 했다.

AI 에이전트 시대에는 다르다. 에이전트가 고객관리 시스템에서 고객의 요구를 읽고, 재무 시스템에서 예산을 확인하고, 마케팅 시스템에서 광고 자리가 비어 있는지 파악해서, 최적의 광고 계획을 순식간에 만들어낸다. 부서의 벽을 넘는 게 아니라, 벽 자체가 존재하지 않는 것처럼 움직인다.

이건 회사가 원해서가 아니다. AI 에이전트가 제대로 일하려면 여러 시스템의 데이터를 동시에 봐야 하기 때문에, 부서의 벽이 있으면 에이전트가 제 기능을 못한다. 벽을 없애는 것은 선택이 아니라 필수 조건이다.

레고 블록 회사와 문어 팀

피라미드도 무너지고, 부서 벽도 허물어지면, 회사는 어떤 모양이 될까.

앞서가는 기업들은 '레고 블록' 같은 조직으로 바뀌고 있다. 레고 블록은 언제든 떼었다 붙일 수 있다. 결제 기능 블록, 고객 데이터 블록, 마케팅 블록. 각각이 독립적으로 작동하면서도, 필요하면 순식간에 조합해서 새로운 형태를 만들 수 있다.

사람 조직도 마찬가지다. 옛날에는 '영업 2팀', '기획 1팀'처럼 고정된 이름의 팀에 배치되면 몇 년간 그 팀에서 일했다. 하지만 새로운 방식에서는, 프로젝트가 생기면 필요한 사람들이 모여서 팀을 만들고, 프로젝트가 끝나면 흩어져서 다른 팀에 합류한다. 이런 유연한 팀을 '모듈형 인재 포드(Modular Talent Pod)'라고 부른다.

로마 공화정 시절의 군대 편성을 떠올리면 이해가 쉽다. 로마는 전쟁마다 시민 중에서 병사를 뽑아 군단을 편성하고, 전쟁이 끝나면 해산했다. 상비군이 아니라 필요할 때 모이는 시민군이었다. 목적에 맞게 모였다가 흩어지는 이 유연함이 로마의 초기 확장에 큰 힘이 되었다.

2026년의 회사도 비슷하게 움직인다. 새 제품을 출시해야 하면, 개발자, 마케터, 디자이너, AI 엔지니어가 하나의 포드로 모인다. 출시가 끝나면 해산하고, 또 다른 프로젝트를 위해 새 포드가 결성된다. 이 포드들 사이에서 AI 에이전트는 서류 작업, 일정 조율, 데이터 정리 같은 뒷일을 맡는다. 사람은 전략과 창의적 판단에 집중한다.

엔지니어에서 정원사로 바뀌는 리더

이런 조직에서 리더의 역할이 완전히 달라진다.

옛날의 리더는 '기계의 설계사'였다. 모든 톱니바퀴가 어떻게 맞물려 돌아가는지를 설계하고, 하나라도 어긋나면 직접 고쳤다. 모든 것을 통제하는 것이 리더의 능력이었다.

하지만 살아있는 유기체 같은 조직에서는 리더가 모든 것을 통제할 수 없다. 세포 하나하나의 움직임을 일일이 지시할 수 없듯이. 대신 리더는 '정원사'가 된다. 정원사는 꽃 한 송이 한 송이에게 "넌 이 쪽으로 피어라, 넌 저쪽으로 피어라"고 명령하지 않는다. 대신 좋은 흙을 갈아주고, 물을 주고, 햇빛이 잘 드는 환경을 만들어준다. 꽃이 스스로 피어나게 한다.

나델라가 마이크로소프트 임원들에게 보낸 내부 메시지에서 "우리 모두 각자의 조직에서 실무자처럼 일하고 행동해야 합니다. 끊임없이 배우고, 배운 것을 버리면서"라고 쓴 것도 같은 맥락이다. 리더가 뒤에서 보고서만 받는 사람이 아니라, 직접 흙을 만지는 정원사가 되라는 뜻이다.

나델라는 한 가지 더 경고했다. "큰 조직이 변화의 속도를 따라가지 못하면, 작은 회사가 이 도구들 덕분에 규모를 얻어서 여러분을 학교에 보내버릴 겁니다." 새로 시작하는 작은 회사는 처음부터 AI에 맞게 조직을 설계할 수 있다. 하지만 큰 회사는 이미 만들어진 피라미드를 허물어야 한다. 그 차이가 승패를 가를 수 있다.

3.2 사람과 여러 AI가 한 팀으로 일하는 법

하나의 완벽한 AI는 없다

회사가 AI를 도입할 때 가장 흔히 저지르는 실수가 있다. "만능 AI 하나만 있으면 다 해결된다"는 생각이다.

현실에서는 그렇게 되지 않는다. 아무리 똑똑한 사람이라도, 회사의 모든 부서 일을 혼자 할 수 없다. 시장 조사도 하고, 보고서도 쓰고, 회계도 처리하고, 고객 응대도 하라고 하면, 그 사람은 과부하에 걸려서 실수를 하거나 쓰러진다.

AI도 마찬가지다. 하나의 AI에게 모든 것을 맡기면, 처리해야 할 일이 너무 많아서 엉뚱한 대답을 지어내거나(이걸 환각, Hallucination이라고 부른다) 중요한 것을 빠뜨린다.

그래서 앞서가는 회사들은 여러 개의 AI를 역할별로 나눠서 쓰고 있다. 이걸 '다중 에이전트 시스템 (Multi-Agent System)'이라고 부른다. 이 시스템은 사람 조직과 놀라울 만큼 닮았다.

데이터브릭스(Databricks)의 2026년 조사에 따르면, 다중 에이전트 시스템의 사용이 불과 4개월 만에 327%나 증가했다. 이건 유행이 아니라 방향이다.

구체적으로 어떻게 돌아가는지 예를 들어보겠다. 수백만 달러짜리 경영 보고서를 만든다고 하자.

총감독 에이전트가 있다. 이 AI는 지휘자 역할을 한다. 전체 작업을 기획하고, 아래의 에이전트들에게 각각 임무를 나눠준다.

정보 수집 에이전트가 있다. 이 AI는 인터넷, 금융 데이터, 뉴스를 돌아다니며 최신 자료를 모아온다.

구조화 에이전트가 있다. 모아온 자료를 논리적으로 정리해서 보고서의 뼈대를 잡는다. 빠진 부분은 없는지, 겹치는 부분은 없는지 확인한다.

작성 에이전트가 있다. 뼈대를 바탕으로 실제 글을 쓰고, 그래프를 그리고, 출처를 달아서 보고서를 완성한다.

검토 에이전트가 있다. 완성된 보고서를 처음부터 끝까지 검사한다. 논리에 구멍이 있으면 돌려보내서 다시 쓰게 한다. 사람에게 넘기기 전에 자체적으로 걸러내는 역할이다.

이 다섯 에이전트가 서로 결과물을 주고받으면서 교차 검증을 한다. 한 에이전트가 실수해도 다른 에이전트가 잡아낸다. 마치 회사의 부서들이 서로 견제하면서 보고서의 품질을 높이는 것과 같다.

로마의 원로원이 왜 독재관(Dictator)보다 나았을까. 한 사람에게 모든 권한을 모으면 빠르지만 위험했다. 여러 사람이 토론하고 견제하면 느리지만 실수가 줄었다. 다중 에이전트 시스템은 이 원리를 AI에 적용한 것이다. 속도와 정확성을 동시에 잡는 방법.

AI의 행동을 통제하는 규칙서와 열쇠

여러 AI가 자율적으로 일하게 하면, 걱정이 생긴다. “제멋대로 행동하면 어찌지?” 당연한 걱정이다. 로마도 총독에게 큰 권한을 줬지만, 동시에 로마법이라는 규칙과 원로원이라는 감시 체계를 갖추고 있었다. 권한과 통제는 항상 짝이다.

AI 에이전트의 행동을 통제하는 장치가 두 가지 있다.

하나는 스킬 파일(Skill.md)이다. 2장에서도 나왔지만, 이건 AI를 위한 업무 규칙서다. “모르는 것을 지어내지 마라. 확인되지 않은 내용은 사람에게 물어봐라. 보고서는 이 양식에 맞춰서 써라.” 이런 규칙이 적혀 있다. 에이전트는 매번 이 규칙서를 읽고 그대로 따른다.

다른 하나는 MCP(Model Context Protocol)다. 이건 에이전트가 회사의 여러 시스템에 접속할 때 쓰는 표준 규격이다. 아무 시스템이나 마음대로 들어가는 게 아니라, “이 에이전트는 캘린더와 슬랙(Slack)만 볼 수 있다. 재무 데이터에는 접근할 수 없다”는 식으로 허락된 범위 안에서만 움직이게 한다.

로마법과 도로의 비유를 다시 쓰자면, 스킬 파일은 법이고, MCP는 통행증이다.

사람이 반드시 끼어들어야 하는 자리

AI 시스템이 아무리 정교해져도, 사람이 빠져서는 안 되는 자리가 있다.

AI는 패턴을 인식하는 데 뛰어나다. 수만 건의 데이터에서 규칙을 찾아내고, 미래를 예측하는 데 사람보다 빠르고 정확할 수 있다. 하지만 AI에게는 양심이 없다. “이 결정이 누군가에게 억울한 피해를 주는 건 아닐까?”를 스스로 느끼지 못한다.

호주에서 실제로 벌어진 일이 있다. 정부가 ‘로보 데트(Robo Debt)’라는 AI 자동화 시스템을 도입해서 복지금을 잘못 받은 사람들에게 돈을 돌려내라는 통지서를 보냈다. 80만 통이 발송됐는데, 그중 상당수가 잘못된 것이었다. 시스템이 계산을 틀렸거나, 사정을 고려하지 않았다. 복지금을 받던 취약 계층의 사람들이 갑자기 빚쟁이가 되었다. 이 사건은 호주 역사상 가장 큰 행정 참사 중 하나로 기록되었다.

속도와 효율에만 눈이 멀어서, “이 결정이 사람에게 어떤 영향을 미칠까”를 생각하는 과정을 빼버렸을 때 벌어지는 일이다.

그래서 ‘사람이 고리 안에 있어야 한다(Human-in-the-loop)’는 원칙이 중요하다.

다만, 사람이 끼어드는 방식도 똑똑해야 한다. AI가 하는 모든 일에 매번 사람이 확인 도장을 찍는 건 답이 아니다. 그러면 AI를 쓰는 의미가 없어진다. 2장에서 말했듯이, 매번 확인하는 ‘절차적 안전장치’ 대신에, 시스템 설계에 안전장치를 심어놓는 ‘구조적 안전장치’가 필요하다.

은행의 대출 심사를 예로 들어보겠다. AI가 대출 신청을 분석할 때, 결과를 세 구역으로 나눈다.

초록 구역. AI가 “이건 확실히 승인해도 된다”고 95% 이상 확신하는 경우. AI가 자동으로 처리한다.

빨강 구역. “이건 확실히 거절 사유가 있다”고 판단되는 경우. 역시 AI가 자동으로 처리한다.

노랑 구역. AI가 확신하지 못하는 경우. 예를 들어 서류 중 뭔가 이상한 점이 있거나, AI가 한 번도 본 적 없는 특이한 상황(예를 들어 생후 9개월 아기 이름으로 된 대출 신청)인 경우. 이때만 사람 전문가에게 알림이 간다.

이렇게 하면 사람은 전체 건수의 대부분을 일일이 들여다보지 않아도 되고, 정말 사람의 판단이 필요한 5%의 어려운 경우에만 집중할 수 있다. AI의 속도와 사람의 판단력을 동시에 살리는 방법이다.

IDC 연구소가 마이크로소프트의 의뢰로 조사한 결과에 따르면, 사람이 의미 있게 참여하는 방식으로 AI를 운영하는 기업은, 그렇지 않은 기업보다 투자 수익이 약 3배 높았다.

켄타우로스가 되는 법

이 모든 것이 합쳐지면, 직장에서 일하는 방식이 근본적으로 달라진다.

직원은 더 이상 도구를 쓰는 ‘작업자’가 아니다. 여러 개의 AI 에이전트를 지휘하는 ‘에이전트 매니저’가 된다. 1장에서 나온 오케스트라 지휘자 비유를 기억하는가. 바이올린을 직접 켜지 않지만, 모든 악기를 하나의 음악으로 만드는 사람.

이 시대에 가장 값어치 있는 인재는, 그리스 신화의 켄타우로스(반인마)를 닮은 사람이다. 사람의 머리와 기계의 몸을 가진 존재. 회사의 사업이 어떻게 돌아가는지를 이해하면서, 동시에 어떤 AI에게 어떤 일을 맡길지를 설계할 수 있는 사람.

중요한 것이 하나 더 있다. 진짜 뛰어난 사람은, 가끔 AI의 스위치를 끌 줄도 안다. AI에게 모든 것을 의존하다 보면, 스스로 깊이 생각하는 능력이 둔해질 수 있다. 좋은 지휘자가 가끔 악보 없이 스스로 멜로디를 흥얼거려보듯이, AI를 쓰는 사람도 가끔은 AI 없이 생각해보는 시간이 필요하다.

카이사르는 뛰어난 장군들을 거느렸지만, 결정적인 순간에는 항상 자기 머리로 판단했다. 갈리아 전쟁에서 알레시아(Alesia) 공성전을 벌일 때, 모든 참모가 후퇴를 건의했지만 카이사르는 이중 포위망을 선택했다. 남의 의견을 듣되, 최종 판단은 자기가 하는 것. AI 시대의 리더에게도 같은 능력이 요구된다.

기계의 힘을 빌리되, 인간의 판단을 잃지 않는 것. 이것이 딱딱한 기계 같은 조직에서 살아 숨 쉬는 생물 같은 조직으로 넘어가는 길이고, 2026년 이후의 세계에서 살아남는 법이다.

KIMKJ.COM

제4장. 미래 직장의 새로운 인재와 직업의 재탄생

4.1 사라지는 직업, 새로 태어나는 직업

가로등 점등원은 어디로 갔을까

19세기 런던의 거리에는 '가로등 점등원'이라는 직업이 있었다. 해 질 무렵이면 긴 막대기를 들고 거리를 걸으며, 가스등에 하나씩 불을 붙이는 사람이었다. 밤의 런던을 밝히는 사람. 당시에는 꽤 중요한 직업이었다.

전기가 발명되자, 이 직업은 사라졌다. 스위치 하나면 수천 개의 전등이 동시에 켜졌다. 점등원의 기술과 경험은 하루아침에 쓸모없어졌다.

하지만 전기는 가로등 점등원을 없앤 대신, 전기 기사, 발전소 운영자, 전선 설계사라는 새로운 직업을 만들어냈다. 사라지는 일자리가 있으면, 생겨나는 일자리도 있었다. 증기기관이 방직 노동자의 60%를 대체했을 때도 같은 일이 벌어졌다. 기계를 다루는 새로운 기술자가 필요해졌다.

2026년, AI라는 새로운 전기가 세상에 퍼지고 있다. 그리고 똑같은 패턴이, 이번에는 훨씬 더 빠른 속도로, 훨씬 더 넓은 범위에서 일어나고 있다.

앤스로픽(Anthropic)의 CEO 다리오 아모데이(Dario Amodei)는 말했다. "앞으로 1년에서 5년 내에 사무직 일자리의 절반이 사라질 것이다." 세계경제포럼(WEF)은 2025년 1월 보고서에서 2030년까지 전 세계에서 9,200만 개의 일자리가 AI에 의해 대체될 것이라고 내다봤다.

숫자만 보면 무섭다. 하지만 같은 보고서에서 WEF는 이렇게도 말했다. 같은 기간 동안 1억 7,000만 개의 새로운 일자리가 생겨난다고. 빼기 9,200만, 더하기 1억 7,000만. 순수하게 7,800만 개의 일자리가 늘어나는 셈이다.

사라지는 것과 생겨나는 것. 이 두 가지를 동시에 봐야 한다. 한쪽만 보면 공포이고, 다른 한쪽만 보면 환상이다. 현실은 그 사이에 있다.

과거의 기술 혁명과 이번의 차이

다만, 이번에는 과거와 다른 점이 하나 있다.

증기기관은 사람의 근육을 대체했다. 무거운 것을 들고, 나르고, 깎는 일. 전기도 마찬가지로 사람의 체력이 필요한 일을 기계에 넘겼다. 하지만 AI는 사람의 머리를 대체하고 있다. 정보를 모으고, 요약하고, 보고서를 쓰고, 코드를 짜고, 데이터를 분석하는 일. 이전까지 "기계가 이걸 못하지"라고 여겨졌던 '생각하는 일'을 AI가 해내기 시작한 것이다.

이 때문에 특히 위험한 직업이 있다. 규칙이 명확하고, 같은 일을 매일 반복하는 사무직이다. 데이터를 입력하는 일, 회의록을 정리하는 일, 기본적인 고객 문의에 답하는 일, 간단한 일정을 조율하는 일. 이런 일은 AI 에이전트가 사람보다 빠르고, 쉬지 않고, 실수 없이 해낸다.

가장 걱정되는 현상은 '속이 빈 중간층'이다. 회사들이 비용을 줄이기 위해 신입 사원이 하던 기초 업무를 AI로 대체하면, 신입 사원이 일을 배울 기회 자체가 사라진다. 초보 시절에 수많은 실수를 하고,

고치고, 다시 해보면서 실력을 쌓는 과정이 없어지는 것이다.

이건 심각한 문제다. 10년 뒤에 이 거대한 AI 시스템을 감독하고 판단할 수 있는 경험 많은 사람이 필요한데, 그런 사람을 키울 훈련장이 사라지고 있으니까. 로마 군단이 신병 훈련소를 없앤다면, 당장은 돈이 절약되겠지만 10년 뒤에는 전쟁을 지휘할 백인대장이 없어진다.

그래서 앞서가는 회사들은 신입 사원의 일을 없애는 대신, 바꾸고 있다. 데이터를 입력하는 일 대신, AI를 이용해 데이터를 분석하고 전략을 제안하는 일을 시키는 것이다. 반복 작업은 AI에게 맡기고, 신입 사원에게는 더 높은 수준의 생각하는 일을 일찍부터 훈련시킨다.

캠브리아기 대폭발처럼 쏟아지는 새 직업들

5억 4천만 년 전, 지구에서 '캠브리아기 대폭발'이라는 사건이 일어났다. 아주 짧은 기간에, 이전에는 존재하지 않았던 수만 종의 새로운 생물이 한꺼번에 쏟아져 나왔다. 과학자들은 왜 이런 일이 벌어졌는지 아직도 연구 중이다.

AI가 일으키고 있는 직업의 변화도 비슷하다. 5년 전에는 이름조차 없었던 직업들이 갑자기 생겨나고 있다.

하나씩 살펴보겠다.

새 직업 1: AI 제품 매니저(AI PM)

2장과 3장에서도 등장했던 이 직업이 왜 이렇게 중요한지, 좀 더 깊이 들어가보겠다.

전 세계 기업의 90%가 AI를 도입했지만, 실제로 돈을 버는 데 성공한 기업은 5%에 불과하다. 95%의 AI 시범 프로젝트가 실패한다는 보고도 있다. 기술이 모자라서가 아니다. 사업하는 사람과 기술 만드는 사람 사이에 말이 안 통하기 때문이다.

AI 제품 매니저는 이 두 세계를 잇는 다리다. 이 사람은 세 가지를 한다.

첫째, 회사 안에서 가장 시간이 많이 들고, AI를 넣었을 때 효과가 큰 지점을 찾는다. 비유하자면, 전쟁에서 적의 가장 약한 곳을 찾아내는 정찰병이다.

둘째, AI가 무엇을 잘하고 무엇을 못하는지 안다. "이 문제에는 글을 쓰는 AI가 필요하고, 저 문제에는 숫자를 예측하는 AI가 필요하다"를 구분한다.

셋째, 사업 부서의 모호한 바람("고객이 덜 떠나게 해주세요")을 기술팀이 바로 일할 수 있는 구체적인 설계도로 바꿔준다.

이 사람은 코딩을 직접 하지 않아도 된다. 대신, 사업의 말과 기술의 말을 동시에 할 줄 알아야 한다.

2026년 기준으로, 경험 있는 AI 제품 매니저의 평균 연봉은 약 16만 달러, 우리 돈으로 2억 원을 넘는다. 기술 업계에서 가장 몸값이 비싼 직업 중 하나가 되었다.

로마 제국에서도 비슷한 역할이 있었다. 원로원의 뜻을 군단장에게 전달하고, 군단의 현장 상황을 원로원에 보고하며, 양쪽의 언어를 모두 이해하는 사절(使節). 전쟁의 승패는 종종 이 사절의 번역 능

력에 달려 있었다.

새 직업 2: AI 윤리 엔지니어

2025년까지의 AI는 대부분 화면 안에서 대화만 나누는 챗봇이었다. 잘못된 답을 해도, 큰 사고로 이어지는 경우는 드물었다.

하지만 2026년의 에이전트 AI는 다르다. 회사의 데이터베이스를 뒤지고, 코드를 실행하고, 고객에게 이메일을 보내고, 결제까지 처리한다. 실제로 행동하는 AI다. 이런 AI가 실수하면, 그 피해는 화면 안에 머물지 않는다. 고객 정보가 유출되거나, 특정 집단을 차별하는 결정이 자동으로 수천 건 실행될 수 있다.

AI 윤리 엔지니어는 이런 사고를 막는 사람이다. 도덕을 논하는 철학자가 아니라, 시스템을 설계하는 기술자에 가깝다.

이 사람이 하는 일은 두 가지다. 하나는 AI가 접근할 수 있는 정보의 범위를 미리 제한하고, 회사의 비밀이 외부로 새어나가지 않게 보안 장치를 만드는 것이다. 다른 하나는 AI의 판단 과정을 추적할 수 있게 기록을 남기고, AI가 공정하게 결정을 내리고 있는지 감시하는 것이다.

로마에 '감찰관(Censor)'이라는 직책이 있었다. 시민의 도덕과 재산을 감시하고, 원로원 의원 중에서 자격 미달인 사람을 퇴출시키는 역할이었다. 권력이 커질수록 감시자도 필요해진다. AI의 권한이 커지는 지금, AI 윤리 엔지니어는 디지털 시대의 감찰관인 셈이다.

새 직업 3: AI 결과물 평가관과 프롬프트 설계자

AI가 만들어낸 보고서, 코드, 분석 결과를 그대로 믿어도 될까. 대답은 "아니오"다.

AI는 95%의 완성도로 그럴듯한 결과물을 만들어낸다. 하지만 나머지 5%에 치명적인 실수가 숨어 있을 수 있다. 없는 사실을 지어내거나(환각), 맥락을 잘못 파악하거나, 미묘한 뉘앙스를 놓치거나. 이 5%를 잡아내는 것이 사람의 몫이다.

AI 결과물 평가관은 AI가 만든 것을 꼼꼼히 검사하는 사람이다. "이 숫자의 출처가 진짜 맞나?", "이 논리에 구멍은 없나?", "이 표현이 고객에게 오해를 살 수 있지 않나?"를 따진다. AI가 95%를 만들고, 사람이 5%의 뉘앙스와 직관을 더해서 100%를 완성하는 구조다.

프롬프트 설계자는 AI에게 일을 시키는 '지시문'을 전문적으로 만드는 사람이다. 단순히 챗GPT에 질문을 잘 던지는 수준이 아니다. 여러 AI 에이전트가 회사의 규칙(스킬 파일 등)에 맞게 정확히 움직이도록 상세한 지침을 설계한다. 에이전트의 '업무 매뉴얼'을 쓰는 사람이라고 보면 된다.

새 직업 4: 다중 에이전트 지휘자

3장에서 다뤘던 '오케스트레이터'가 직업으로 자리를 잡고 있다.

인사 부서의 채용 담당자를 예로 들어보겠다. 예전에는 이력서를 직접 읽고, 면접 일정을 잡고, 합격 통지서를 쓰는 것이 이 사람의 일이었다. 이제는 이력서를 선별하는 AI, 면접 일정을 잡는 AI, 합격 통지를 보내는 AI를 배치하고 지휘하는 것이 이 사람의 일이 된다. 직접 하는 게 아니라, 시키는 것이다.

이런 역할은 모든 부서에서 생겨나고 있다. 마케팅, 영업, 재무, 고객 지원. 어디서든 여러 개의 AI 에이전트를 조합해서 하나의 팀처럼 운영하는 사람이 필요해졌다.

이 모든 새 직업에는 공통점이 하나 있다. "기계보다 계산을 빨리 하는 사람"이 아니라, "기계에게 무엇을, 왜 시킬지를 정하는 사람"이라는 것이다.

4.2 슈퍼 개인의 등장: 한 명이 수십 명의 일을 해내는 시대

사람을 더 뽑으면 매출이 느는 시대는 끝났다

로마 제국이 확장할 때, 더 넓은 영토를 다스리려면 더 많은 병사와 관리가 필요했다. 군단을 하나 더 편성하고, 총독을 한 명 더 보내야 했다. 영토의 크기와 사람 수가 비례했다.

2020년대 초반까지, 회사도 같은 방식으로 성장했다. 고객이 늘면 직원도 늘려야 했다. 영업 사원을 더 뽑고, 마케터를 더 채용하고, 관리자를 더 앉혔다. 매출과 직원 수가 함께 올라가는 것이 당연한 공식이었다.

AI 에이전트가 이 공식을 깨뜨렸다.

이제 한 사람이 AI 에이전트 여러 개를 부리면서, 예전에 여러 사람이 하던 일을 혼자서 해낼 수 있게 되었다. 회사의 매출은 올라가는데, 직원 수는 그대로이거나 오히려 줄어드는 현상이 벌어지고 있다.

이런 사람을 '슈퍼 개인(Super IC, Super Individual Contributor)'이라고 부른다. IC는 '개인 기여자', 그러니까 팀을 관리하는 게 아니라 직접 일하는 사람이라는 뜻이다. 거기에 '슈퍼'가 붙었다. AI의 힘을 빌려서 한 사람이 열 명, 스무 명 몫의 결과를 내는 사람.

500달러로 5만 달러짜리 팀을 만든 사람

제이콥 बैं크(Jacob Bank)라는 사람이 있다. 구글에서 지메일(Gmail)과 캘린더의 제품 책임자로 일하다가 나와서, 릴레이(Relay.app)라는 AI 에이전트 회사를 창업했다. 2026년 3월에 화제가 된 인터뷰에서 그의 이야기가 나왔다.

뱅크는 이 회사에서 유일한 마케팅 담당자다. 마케팅팀이라고 할 것도 없다. 그 혼자니까. 그런데 그의 지휘 아래에는 약 40개의 AI 에이전트가 일하고 있다.

어떤 에이전트는 경쟁사의 움직임을 24시간 추적한다. 경쟁사가 가격을 바꾸면, 즉시 슬랙(Slack)으로 알림이 온다. 어떤 에이전트는 매주 업계 트렌드를 분석해서 보고서를 만든다. 어떤 에이전트는 그 트렌드를 바탕으로 링크드인(LinkedIn) 게시글 초안을 쓴다. 어떤 에이전트는 웨비나(온라인 세미나) 일정이 잡히면 자동으로 홍보 이메일을 보낸다.

이 40개 에이전트를 운영하는 데 드는 비용은 한 달에 500달러, 우리 돈으로 약 65만 원이다.

만약 같은 일을 사람에게 시키려면, 최소 4명의 전문 마케터가 필요하다. 한 달 인건비로 5만 달러, 약 6,500만 원이 든다. 비용 차이가 100배다.

뱅크는 이 시스템으로 링크드인에서 150만 건의 노출을 기록했다.

한 가지 더 흥미로운 점이 있다. 뱅크는 고객에게 회의록을 예쁘게 정리해서 보내는 에이전트를 만들었다가, 고객이 그걸 안 읽는다는 것을 알아차렸다. 그는 즉시 그 에이전트를 '해고'했다. 사람 직원을 해고하려면 눈치도 봐야 하고, 감정적 부담도 크고, 절차도 복잡하다. AI 에이전트는 쓸모없으면 끄면 된다. 전략이 바뀌면 "하던 일 멈춰"라고 한마디 하면 그만이다.

뱅크 본인의 말이 인상적이다. "만약 내 부모님 세대였다면, 대기업에서 40년 일하는 게 안전한 길이었을 겁니다. 하지만 지금은 반대입니다. 한 회사에만 너무 오래 머물면서 그 회사에서만 통하는 기술만 쌓는 게 가장 위험한 길입니다."

7명이 750억 원을 벌어들인 회사

대만에 '대화국제(戴華國際)'라는 메모리 반도체 수출 회사가 있다. 이 회사의 핵심 인력은 단 7명이다. 그런데 1년 매출이 18억 대만달러, 우리 돈으로 약 750억 원이다.

비결은 이렇다. 이 7명은 10년 넘게 글로벌 시장을 돌아다니며 쌓은 경험과 판단력을 AI 시스템에 녹여넣었다. 'AO 무역통'이라는 자체 AI를 만들어서, 전에는 전문 인력이 18시간 걸려 쓰던 마케팅 문서를 15분 만에 완성하게 했다. 72배 빨라진 것이다.

이 사람들은 비싼 해외 박람회에 가는 대신, AI가 24시간 전 세계 고객을 상대하게 만들었다. 사람은 자고 있어도, AI는 지구 반대편의 고객에게 이메일을 보내고, 문의에 답하고, 제안서를 작성한다.

7명이 750억 원을 버는 것이 가능한 이유는, 이 7명이 단순한 작업자가 아니기 때문이다. 10년의 경험에서 나온 판단력과 통찰을 AI에 주입하고, AI가 그 통찰을 수천 배로 복제해서 실행하는 구조다. 사람의 지혜를 AI의 규모로 증폭시킨 것이다.

코딩 도구 커서(Cursor)도 비슷한 사례다. 직원이 60명뿐인데, 전통적인 1,000명 규모의 소프트웨어 회사보다 매출이 많다. 2026년 3월 기준 연매출 20억 달러(약 2조 8천억 원)를 돌파했다.

사람 수 대신 AI 사용량에 투자하는 시대

이런 변화를 가리키는 재미있는 표현이 있다. '헤드카운트 극대화(Headcount Maxing)'에서 '토큰 극대화(Token Maxing)'로 바뀌었다는 것이다.

헤드카운트는 직원 수다. 토큰은 AI를 사용할 때 소비되는 단위, 쉽게 말하면 AI 사용료다.

예전에는 회사가 성장하려면 사람을 더 뽑아야 했다. 직원 한 명을 고용하면 월급, 보험료, 사무실 공간, 복지비 등으로 연간 수천만 원이 든다. 그리고 그 사람은 하루 8시간, 주 5일만 일한다.

하지만 AI 에이전트는 24시간, 365일 일한다. 비용은 해마다 급격히 떨어지고 있다. 그래서 뛰어난 회사들은 사람을 한 명 더 뽑는 대신, 지금 있는 뛰어난 사람에게 AI를 마음껏 쓸 수 있는 예산을 준다. 한 명의 실력자가 AI를 이용해서 열 명 몫의 결과를 내는 것이 더 효율적이기 때문이다.

회사 입장에서 수천만 원짜리 AI 사용료 영수증은 무섭게 보일 수 있다. 하지만 그 AI가 대체하고 있는 것은 수억 원의 인건비와 관리 비용이다. 계산은 명확하다.

슈퍼 개인이 아침에 출근하면

슈퍼 개인의 하루가 어떻게 돌아가는지 구체적으로 그려보겠다.

아침에 출근한다. 밤사이 AI 정보 수집 에이전트가 50개 경쟁사의 움직임을 요약해놓았다. 읽어보니, 경쟁사 하나가 새 제품을 발표했다. 슈퍼 개인은 AI 작성 에이전트에게 "이 경쟁사의 새 제품에 대한 우리 회사의 대응 전략 초안을 써줘"라고 지시한다. 10분 뒤 초안이 도착한다. 읽어보고, AI가 놓친 부분(고객의 감정적 반응이나, 우리 회사만 아는 내부 사정)을 직접 5% 추가한다. 그리고 승인 버튼을 누른다.

이 사람의 역할은 세 가지다. 방향을 정하는 것. AI가 못 보는 인간적인 부분을 채우는 것. 최종 결과물에 책임을 지는 것.

RPM이라는 틀을 기억하는가. 1장에서 나왔던 것이다.

R(역할). 에이전트 A는 정보 수집을 하는 '정보 장교'. 에이전트 B는 글을 쓰는 '작가'. 에이전트 C는 회사 규칙에 맞게 다듬는 '편집장'.

P(프로젝트). "이번 주까지 경쟁 분석 보고서 완성."

M(방법). "출처 없는 내용은 쓰지 마. 회사 템플릿에 맞춰서 써. 모르면 지어내지 말고 알려줘."

코딩을 몰라도 된다. 사업의 논리와 목표를 명확히 아는 것이 핵심이다. "무엇을, 왜 해야 하는지"를 정하는 능력. 이것이 슈퍼 개인의 무기다.

기보를 설계하는 사람이 이긴다

1장에서 체스 비유를 들었다. 다시 꺼내보겠다.

조만간 모든 사람이 똑같이 강력한 AI라는 기물을 손에 쥐게 된다. 그때 승부를 가르는 것은 기물 자체가 아니라, 기보(棋譜)를 설계하는 능력이다. 어떤 순서로, 어떤 목적을 가지고, 어디에 배치하느냐.

수십 명이 이메일과 엑셀을 주고받으며 낭비하던 시간을 AI로 압축하고, 오직 사람만이 할 수 있는 판단과 창의에 에너지를 쏟아붓는 사람. 이 사람이 슈퍼 개인이다.

카이사르가 쓴 '갈리아 전쟁기'를 읽으면, 그가 전투에서 직접 칼을 휘두르는 장면은 거의 나오지 않는다. 대신, 어느 군단을 어디에 보낼지, 보급선을 어떻게 유지할지, 갈리아 부족 사이의 갈등을 어떻게 이용할지를 끊임없이 설계하는 모습이 나온다. 직접 싸우는 것은 병사에게 맡기고, 전쟁의 그림을 그리는 것이 카이사르의 일이었다.

2026년의 직장도 같은 구조로 가고 있다. 실행은 AI에게 맡기고, 전략을 설계하는 것이 사람의 일이 된다.

가로등 점등원은 사라졌지만, 도시의 밤은 더 밝아졌다. 사라지는 직업을 아쉬워할 필요는 없다. 다만, 새로 생겨나는 직업이 요구하는 능력을 지금부터 준비해야 한다. 그 능력의 핵심은 단순하다. 기계에게 시킬 줄 아는 것, 그리고 기계가 못하는 것을 아는 것.

KIMKJ.COM

제5장. 성공을 이끄는 리더의 자세와 조직 문화 바꾸기

5.1 사장이 먼저 운전대를 잡아야 한다

슈퍼카를 사놓고 대리 운전을 시키는 사장들

율리우스 카이사르가 갈리아 원정을 떠날 때, 그는 로마에 앉아서 보고서만 기다리지 않았다. 직접 말에 올라타고, 직접 전장을 걸었다. 갈리아의 땅이 어떤 냄새가 나는지, 병사들의 사기가 어떤지, 적의 진형이 실제로 어떻게 보이는지를 자기 눈으로 확인했다. 보고서에 적힌 갈리아와, 직접 밟아본 갈리아는 전혀 다른 땅이었다.

2026년, AI라는 새로운 전장에서 대부분의 기업 리더들은 카이사르와 정반대로 행동하고 있다.

전 구글 대만 총괄 책임자였던 젠리펑(簡立峰) 박사가 뼈아픈 비유를 했다. 수십억 원을 들여 최첨단 AI라는 슈퍼카를 사놓고, 사장 본인은 운전대를 잡지 않는다는 것이다. 두렵거나 귀찮기 때문이다. 대신 평생 자전거만 타본 비서에게 운전을 맡기고, 자기는 뒷좌석에 앉아서 “왜 더 빨리 안 달리냐”고 다그린다. 대만 중소기업 경영진의 80%가 AI를 직접 쓰지 않고 비서에게 떠넘긴다는 조사 결과가 나왔는데, 이건 대만만의 문제가 아니다.

리더가 뒷좌석에 앉아 있으면, 회사 전체가 자전거 속도로 달린다. 아무리 좋은 엔진을 사도 소용이 없다.

PwC의 조사에 따르면, AI를 도입한 기업 중 56%가 아무런 성과를 보지 못했다. AI 시범 프로젝트의 95%가 실패했다는 보고도 있다. 기술의 문제가 아니다. 리더가 직접 기술을 만져보지 않아서, 그 기술이 진짜로 무엇을 할 수 있고 무엇을 못하는지 모르기 때문이다.

직접 뛰는 사장: ICCEO라는 새로운 리더

옛날 경영학 교과서에는 이렇게 적혀 있었다. “좋은 CEO란 세부적인 실무에서 손을 떼고, 권한을 나눠주고, 높은 곳에서 큰 그림만 보는 사람이다.” 하지만 AI 시대에는 이 공식이 틀렸다.

에어테이블(Airtable)이라는 기업 가치 약 120억 달러(약 17조 원)의 회사를 만든 하위 류(Howie Liu) CEO는 완전히 다른 리더의 모습을 보여준다. 그는 스스로를 ‘ICCEO(Individual Contributor CEO)’, 우리말로 하면 ‘직접 뛰는 최고경영자’라고 부른다.

하위 류는 자기 회사에서 AI 사용 비용(추론 비용)을 가장 많이 쓰는 사람이다. 1위다. 직접 코드를 짜고, 회사에서 지난 1년간 녹음된 수만 건의 영업 회의를 AI로 분석한다. API 호출 비용으로 수백 달러를 기꺼이 쓴다. 그 결과, 외부 컨설팅 회사에 수백만 달러를 주고 몇 달을 기다려야 알 수 있었던 것들(고객이 진짜로 원하는 것, 가장 효과적인 영업 방법, 시장에서 우리 회사의 위치)을 며칠 만에 직접 파악해낸다.

2026년 초, 하위 류는 더 파격적인 행보를 보였다. 정기적으로 잡혀 있던 임원 회의 대부분을 취소하고, 대신 팀원들에게 “일주일 동안 모든 회의를 취소하고, 시장에 나와 있는 모든 AI 제품을 직접 써 봐라”고 권했다. 직접 만져보고, 부딪혀보고, 실패해보라는 뜻이다.

그는 이렇게 말한다. CEO가 회사의 '수석 미각 책임자(Chief Taste Officer)'가 되어야 한다고. 요리사가 외계에서 온 새로운 식재료를 얻었을 때, 설명서만 읽고 밑의 요리사에게 시키면 그 식재료의 진짜 가치를 알 수 없다. 총주방장이 직접 만져보고, 냄새 맡고, 씹어봐야 한다.

실제로 하위 류는 자기가 매일 클로드(Claude)를 직접 쓰는 '헤비 유저'라고 공개적으로 밝혔다. 에어테이블은 2026년 1월에 '슈퍼에이전트(Superagent)'라는 새 AI 제품을 내놓았고, 4월에는 '하이퍼에이전트(Hyperagent)'까지 공개했다. 13년 역사상 처음으로 본업 외의 새 제품을 내놓은 것인데, 이것이 가능했던 이유는 CEO가 직접 AI의 가능성을 몸으로 느꼈기 때문이다.

리더의 자기 혁명 4단계

"좋은 이야기다. 하지만 나는 기술을 잘 모른다. 어떻게 시작하란 말인가?" 이런 질문을 할 리더가 많을 것이다.

젠리핑 박사가 제안하는 4단계가 있다.

1단계. 고통을 겪어라.

잔인하게 들리겠지만, 효과는 확실하다. 리더의 비서를 1~2주간 휴가 보내라. 일정 조율, 이메일 분류, 회의 자료 준비, 데이터 정리. 이 모든 '잡일'을 리더가 직접 맨몸으로 해보라. 죽을 만큼 힘들 것이다. 바로 그 고통 속에서, "우리 회사의 일이 도대체 어디서 막히고 있는지"를 처음으로 온몸으로 느끼게 된다.

로마의 하드리아누스 황제는 제국 전역을 직접 걸어서 시찰한 것으로 유명하다. 로마에 앉아서 보고서만 읽었다면, 브리타니아(영국)에 성벽이 필요하다는 판단을 내릴 수 없었을 것이다. 직접 걸어보고, 직접 추위를 느끼고, 직접 적의 침입 루트를 눈으로 봤기 때문에 하드리아누스 방벽이라는 역사적 결정을 내릴 수 있었다.

2단계. 도구를 직접 써봐라.

고통을 겪으면, 해결책을 찾고 싶은 절박함이 생긴다. 비서 없이 허덕이는 리더는 살기 위해 ChatGPT나 클로드를 직접 켜게 된다. 프레젠테이션 초안을 AI에게 시켜보고, 꼬여있는 이메일에 답장을 써달라고 해보고, 난잡한 회의록을 정리시켜본다. 서투르겠지만, 그 서투름이 배움이다.

3단계. 가능성을 발견하라.

AI를 직접 써보다 보면, 머리를 한 대 맞은 듯한 순간이 온다. "10년 경력 비서가 며칠 걸려 만들어오던 보고서보다, AI가 5분 만에 뽑아낸 초안이 더 낫잖아?" 이 순간, AI가 잡일 처리기가 아니라 최고의 전략 파트너가 될 수 있다는 걸 깨닫는다.

4단계. 술선수범하라.

여기까지 온 리더는 더 이상 외부 컨설턴트의 말에 의지하지 않는다. 회사의 어떤 부서, 어떤 과정에 AI를 넣어야 할지 스스로 안다. 직접 겪어봤으니까. 이 리더가 밀어붙이는 변화는 진짜다. 보고서에서 읽은 변화가 아니라, 몸으로 느낀 변화이기 때문이다.

중소기업에서는 사장이 곧 전략이다

이 이야기가 대기업에만 해당한다고 생각하면 오산이다. 사실 중소기업에서 이 문제가 더 절실하다.

3,000명짜리 대기업에서 AI는 ‘비용 줄이기’ 도구가 될 수 있다. 직원 3,000명이 각각 5%의 시간만 아껴도, 합치면 엄청난 양의 인건비가 절약된다.

하지만 직원 20~30명짜리 중소기업은 다르다. 직원들이 엑셀을 30% 빨리 끝내고 오후 4시에 퇴근한다고 해서, 회사 매출이 늘어나지 않는다. 왜냐하면 중소기업의 매출 한계는 직원의 속도가 아니라, 사장의 역량에 묶여 있기 때문이다.

중소기업에서 AI의 진짜 목표는 직원의 보고서 시간을 5분 줄이는 게 아니다. 사장의 시간을 해방시키는 것이다. 사장이 정보 수집, 이메일 대응, 일정 조율 같은 잡일에 빼앗기는 시간을 AI에게 넘기고, 그 시간으로 핵심 고객을 만나고, 새 사업을 구상하고, 결정적인 계약을 따내야 한다.

나델라가 다보스에서 50개 넘는 양자 회의를 준비하면서 AI를 직접 쓴 이야기를 3장에서 했다. 그는 AI에게 “내일 만날 CEO와의 관계, 투자 상황, 현재 이슈를 360도로 브리핑해줘”라고 지시하고, 몇 초 만에 완벽한 요약물을 받아들였다. 나델라처럼 세계 최대 IT 기업의 수장도 직접 AI를 쓰는데, 중소기업 사장이 비서에게 맡기고 앉아 있을 이유가 있을까.

5.2 보이지 않는 빛, ‘문화 부채’를 갚아야 한다

눈에 안 보이지만 쌓이고 있는 빛

로마 제국이 최전성기를 지나 쇠퇴하기 시작했을 때, 겉으로 보기에는 문제가 없었다. 군대는 여전히 강했고, 도로는 여전히 깔려 있었고, 법은 여전히 작동했다. 하지만 눈에 보이지 않는 곳에서 균열이 시작되고 있었다. 시민들의 로마에 대한 충성심이 흔들리고, 군인들의 사기가 떨어지고, 원로원 의원들이 공공의 이익보다 자기 이익을 앞세우기 시작했다. 이 눈에 보이지 않는 균열이 결국 제국을 무너뜨렸다.

기업에도 이런 눈에 보이지 않는 균열이 있다. 이것을 ‘문화 부채(Culture Debt)’라고 부른다.

빛은 갚지 않으면 이자가 붙듯이, 문화 부채도 방치하면 점점 커진다. 그리고 어느 순간, 회사를 안에서 서부터 무너뜨린다.

문화 부채는 이렇게 쌓인다. 회사가 AI를 도입하면서 오직 효율과 성과에만 신경 쓰고, 직원들의 불안감은 무시한다. 56%의 기업이 AI 도입 때 직원의 감정과 공정성을 완전히 외면했다는 조사 결과가 있다. 딜로이트(Deloitte)에 따르면, 요즘 직원들은 1년에 무려 15번의 큰 조직 변화를 겪어야 한다. 새 시스템, 새 도구, 새 절차가 끊임없이 밀려온다. 70%의 직원이 ‘변화 피로(Change Fatigue)’를 호소한다. 지쳐서 더 이상 따라가기 싫은 것이다.

직원들이 AI 실력을 숨기는 이유

3장에서 ‘새도우 AI’ 이야기를 했다. 직원들이 회사 몰래 개인 AI 도구를 쓰는 현상이다. 여기서 좀 더 깊이 들어가보겠다.

마이크로소프트와 링크드인의 최신 보고서에 따르면, 직원 개인은 AI를 무서운 속도로 배우고 있다. 주말과 퇴근 후 시간을 쪼개서, 혼자 AI 사용법을 익히고, 업무를 10배 빠르게 처리하는 방법을 알아

낸다. 65%의 직원이 “AI를 안 쓰면 뒤처질 것 같다”는 불안을 느낀다.

그런데 55%의 직원은 자기가 AI를 잘 쓰게 된 것을 회사에 숨긴다.

왜 그럴까. 이유가 썩썩하다. 자신이 AI로 놀라운 성과를 내더라도, 회사에서 칭찬이나 보상을 받을 거라고 기대하는 직원은 13%밖에 안 된다. 87%는 “알아주지도 않을 거다”라고 생각한다. 심지어 “AI로 내 업무를 빨리 끝내는 법을 회사에 알려주면, 회사가 날 해고하는 거 아닐까?”라는 두려움도 있다. 내가 가진 노하우를 AI에 넘겨주면, 내가 쓸모없어지는 거 아닐까.

이 두려움은 이해할 만하다. 하지만 결과는 파괴적이다. 78%의 직원이 회사 승인 없이 개인 AI 도구를 쓰고 있고, 38%는 회사의 비밀 정보를 AI에 넣어본 적이 있다. 회사는 직원들이 어디서 뭘 하는지 모른다. 정보 유출의 위험이 커지고 있는데, 아무도 보지 못한다.

로마 제국 말기에, 변방의 군단들이 로마 중앙의 지시를 따르지 않고 자기들끼리 움직이기 시작했던 현상과 비슷하다. 중앙이 현장의 현실을 모르니까, 현장이 중앙을 무시하기 시작한 것이다.

실수해도 괜찮은 회사가 이긴다

이 문제의 열쇠는 ‘심리적 안전감(Psychological Safety)’이다. 2장에서도 나왔지만, 워낙 중요한 개념이라 다시 한번 깊이 다루겠다.

심리적 안전감은 “실수해도 혼나지 않는 분위기”다. “바보같은 질문을 해도 무시당하지 않는 환경”이다. “새로운 걸 시도하다 실패해도, 그 실패 자체가 가치 있다고 인정받는 문화”다.

이건 ‘좋은 직장’의 조건’ 같은 부드러운 이야기가 아니다. 돈과 직결되는 문제다.

직원들이 두려움 없이 새로운 시도를 할 수 있는 팀에서는 새로운 아이디어가 27% 더 많이 나온다. 결정을 내리는 속도가 35% 빨라진다. 핵심 인재가 회사를 떠나는 비율이 50%나 줄어든다.

AI 시대에 이 안전감이 더욱 중요해진 이유가 있다. AI는 완벽하지 않다. 가끔 환각(Hallucination)을 일으키고, 편향된 결과를 내놓는다. “AI를 써봤는데 엉뚱한 결과가 나왔어요”라고 솔직하게 말할 수 있어야 한다. “완벽하게 하지 않으면 혼난다”는 분위기에서는, 아무도 새로운 도구를 시도하지 않는다. 실패가 두려우니까.

사장이 먼저 실패를 보여줘라

심리적 안전감을 만드는 가장 강력한 방법은 무엇일까. 마이크로소프트와 링크드인의 연구에서 답이 나왔다. “리더가 직접 AI를 쓰는 모습을 보여줄 때, 직원들의 AI에 대한 신뢰와 사용이 폭발적으로 늘어난다.”

더 효과적인 것은, 리더가 실패를 공개하는 것이다.

회의 시간에 사장이 이렇게 말한다고 상상해보자. “나도 어제 AI로 프레젠테이션을 만들어봤는데, AI가 엉뚱한 숫자를 지어내더라. 완전히 망했어. 그런데 이 부분은 정말 놀라웠다. 내가 세 시간 걸릴 일을 10분 만에 끝내줬어.”

이 한마디가 직원들에게 강력한 메시지를 보낸다. "사장도 실패하는구나. 나도 실패해도 괜찮겠구나. " 리더의 이런 솔직함이 직원들에게 면죄부를 준다. AI를 시도해봐도 된다는 허락.

나델라가 마이크로소프트 임원들에게 보낸 내부 메시지도 같은 맥락이었다. "우리 모두 실무자처럼 일하고, 끊임없이 배우고, 배운 것을 버려야 합니다." 그리고 "AI 스타트업에서 일하는 친구가 얼마나 다르게 일하는지 이야기하는 메모를 받을 때마다 웃음이 납니다. 사실 그런 혁신은 바로 여기 마이크로소프트에서도 일어나고 있습니다. 우리의 할 일은 그것을 찾아내고, 키워주고, 배우는 것입니다."

"해고당할까 봐 두렵다"에 대한 정면 돌파

심리적 안전감을 갉아먹는 가장 깊은 뿌리는 실직의 두려움이다. "내가 AI에게 내 업무 방법을 다 가르쳐주면, 회사는 나를 자를 텐데?"

대부분의 리더는 이 질문을 피한다. "AI는 여러분을 대체하지 않아요, 보조할 뿐이에요"라는 애매한 위로로 넘어간다. 하지만 직원들은 바보가 아니다. 그 위로를 믿지 않는다.

진짜 신뢰를 쌓으려면 솔직해야 한다. IBM이 보여준 방식이 참고할 만하다. "자동화로 인해 특정 반복 업무는 사라질 것입니다. 이건 사실입니다." 뼈아픈 진실을 먼저 밝힌다. 그리고 동시에 약속한다. "하지만 사람만이 할 수 있는 새로운 직무를 만들 것이고, 여러분이 그 자리로 이동할 수 있도록 교육을 지원하겠습니다."

모호한 위로보다, 아픈 진실과 구체적인 약속이 신뢰를 만든다.

회사 안에 숨어 있는 AI 천재를 찾아라

마지막으로, 가장 흥미로운 이야기를 하겠다.

회사가 AI를 도입하려고 할 때, 많은 리더가 비싼 외부 AI 전문가를 채용하려 한다. 하지만 진짜 보물은 이미 회사 안에 있을 수 있다.

구매팀에서 10년 동안 복잡한 발주 프로세스를 관리해온 사람, 재무팀에서 수백 개의 엑셀 파일을 조합해 보고서를 만들어온 사람, 인사팀에서 채용부터 퇴직까지 전 과정을 꿰뚫고 있는 사람. 이 사람들은 코딩을 모를 수 있다. 하지만 업무의 흐름을 누구보다 잘 안다. 어디가 막히는지, 어디서 시간이 낭비되는지, 어떤 순서로 일해야 효율적인지를.

AI 에이전트 시대에는 코딩 없이도 AI를 조합해서 업무를 자동화할 수 있다. 이 사람들이 AI 도구를 손에 쥐면, 자기 부서의 골칫거리를 스스로 해결하는 시스템을 만들어낼 수 있다. 왜냐하면 문제가 뭔지를 가장 잘 아는 사람이기 때문이다.

리더의 역할은 이 숨겨진 천재들에게 '안전한 놀이터'를 만들어주는 것이다. 보안 규칙은 지키되, 그 안에서는 마음껏 실험해도 되는 공간. 성과를 당장 ROI(투자 수익률)로 따지지 않고, "재밌는 거 발견했어?"라고 물어주는 분위기. 이 사람들이 자기 부서의 문제를 해결하면, 자연스럽게 주변 동료들에게도 퍼진다. "저 사람도 했는데, 나도 해볼까?" 이것이 가장 강력한 내부 전파 방식이다.

로마 제국이 멀리 확장할 수 있었던 이유 중 하나는 '로마 시민권'이라는 보상이었다. 변방의 군인이거나 속주의 주민이라도, 로마에 충성하고 기여하면 시민권을 받았다. 이 시민권은 새로운 기회와 권리를 의미했다. 기여에 대한 명확한 보상이 있으니까, 사람들은 제국을 위해 스스로 뛰었다. 회사에서도 마찬가지다. AI를 활용해 혁신을 만들어낸 직원에게 인정과 보상을 주면, 더 많은 사람이 스스로 움직인다.

배우는 방식도 바뀌어야 한다

세계적인 인재 분석가 조시 버신(Josh Bersin)에 따르면, 전 세계 기업이 매년 직원 교육에 4,000억 달러(약 560조 원)를 쓴다. 하지만 74%의 기업은 "우리 교육이 변화 속도를 못 따라간다"고 인정한다.

왜 그럴까. 옛날 방식의 교육은 '밀어넣기(Push)'다. 회사가 교육 일정을 잡고, 직원을 교실에 앉히고, 정해진 내용을 듣게 한다. 강의를 들으면서 졸리기 시작하고, 끝나면 90%를 잊어버린다.

AI 시대의 교육은 '끌어당기기(Pull)'로 바뀌어야 한다. 직원이 업무를 하다가 막히는 순간, AI가 실시간으로 필요한 정보를 찾아서 보여주는 것이다. 영업 사원이 제안서를 쓰다가 막히면, AI가 그 고객의 업종 트렌드와 과거 성공 사례를 자동으로 띄워준다. 교실에 앉아서 듣는 게 아니라, 일하면서 배우는 것이다.

필요한 순간에, 필요한 것을, 바로 배우는 것. 이것이 AI 시대의 교육이다.

결론은 기술이 아니라 문화다

이 장에서 한 모든 이야기를 한마디로 줄이면 이것이다.

AI 시대에 기업의 승패를 가르는 것은 "어떤 AI를 샀느냐"가 아니라, "직원들이 AI를 두려움 없이 쓸 수 있는 문화를 만들었느냐"다.

리더가 먼저 뛰어들고, 실패를 솔직하게 공유하고, 직원의 두려움에 정면으로 답하고, 숨겨진 인재에게 놀이터를 제공하고, 배우는 방식을 현장 중심으로 바꾸는 것.

로마가 위대했던 이유는 군사력만이 아니었다. 패배한 적에게도 시민권을 주고, 다른 문화를 흡수하고, 변화를 두려워하지 않았던 개방적인 문화가 있었기 때문이다. 로마가 쇠퇴한 것도 군사력이 약해져서가 아니라, 그 개방적 문화가 닫히기 시작했을 때부터였다.

2026년의 기업도 마찬가지다. AI라는 무기는 누구나 살 수 있다. 하지만 그 무기를 제대로 쓸 수 있는 문화를 가진 조직은 드물다. 문화가 곧 해자(壕字)다. 경쟁자가 돈으로 살 수 없는 유일한 방어선이다.

KIMKJ.COM

제6장. 작은 실험을 회사 전체로 넓히는 전략

6.1 시범 사업의 연옥에서 탈출하는 법

성공한 실험이 왜 회사를 바꾸지 못하는가

로마가 갈리아를 정복할 때, 첫 번째 전투에서 이기는 것은 시작에 불과했다. 한 마을을 점령하는 것과, 갈리아 전체를 로마의 속주로 만드는 것은 완전히 다른 차원의 일이었다. 한 마을에서 통한 전술이 다른 마을에서는 안 먹히기도 했고, 현지 부족의 저항이 예상보다 거셌고, 보급선이 늘어나면서 예상치 못한 문제가 터져 나왔다.

회사에서 AI를 도입하는 과정도 똑같다.

작은 팀에서 AI 시범 사업(파일럿)을 해본다. 결과가 좋다. 경영진 앞에서 시연하면 "와, 대단하다"는 반응이 나온다. 그런데 18개월이 지나도, 회사의 일하는 방식은 전혀 바뀌지 않는다. 시범 사업은 성공했는데, 회사 전체로 퍼지지 않는 것이다.

이 상태를 '파일럿의 연옥(Pilot Purgatory)'이라고 부른다. 연옥은 천국도 지옥도 아닌 곳이다. 시범 사업이 성공도 아니고 실패도 아닌 상태로 떠돌아다니는 것.

전 세계 기업의 70% 이상이 AI를 도입하면서 이 연옥을 경험한다. 기술이 안 돼서가 아니다. AI 모델은 잘 작동한다. 문제는 사람과 조직이 변하지 않는 데 있다. 새 기술이 옛날 방식의 일과 부딪히면, 조직은 기술의 힘을 흡수하는 대신 밀어내버린다.

그래서 작은 실험을 회사 전체로 넓히려면, AI를 '기술 사업'이 아니라 '사람과 조직의 변화'로 접근해야 한다. 이때 쓸 수 있는 틀이 SHIFT라는 5단계 방법이다.

S: 이야기와 전략, 왜 하는지를 먼저 말하라

대부분의 회사가 AI를 소개할 때 이런 말부터 한다. "이 새로운 챗봇은 문서를 10배 빨리 요약해줍니다." 기능 설명부터 시작하는 것이다.

하지만 직원들이 궁금한 것은 기능이 아니다. "왜 이걸 하는 거지?" "이게 나한테 무슨 의미가 있지?" "혹시 내 자리가 없어지는 건 아닌가?"

리더가 이 질문에 먼저 답해야 한다. "우리가 이걸 하는 이유는 이것이다. 이것이 회사의 생존에 이렇게 연결된다. 그리고 이것이 여러분 개인에게 이렇게 도움이 된다."

명확한 이야기(Story)가 없으면, 그 빈자리를 두려움이 채운다.

카이사르가 갈리아 원정을 시작할 때, 병사들에게 "우리는 왜 이 전쟁을 하는가"를 분명히 설명했다. 전리품의 약속만이 아니었다. 로마의 안전을 위해, 로마의 영광을 위해. 목적이 분명한 군대는 강하다. 목적 모르고 싸우는 군대는 첫 번째 패배에서 흩어진다.

H: 직원의 감정을 이해하라

AI 도입은 새 프로그램 하나 깔아주는 것과 다르다. 직원의 '정체성'을 건드리기 때문이다.

10년 경력의 베테랑 직원에게 “이제 AI가 당신이 하던 보고서를 쓸 겁니다”라고 말하는 것은, “당신이 10년간 갈고닦은 기술이 이제 기계로 대체됩니다”라는 뜻으로 들린다. 이건 새 도구를 배우라는 요구가 아니다. 자기 존재 가치에 대한 위협이다.

변화를 연구하는 학자 윌리엄 브리지스(William Bridges)는 사람이 변화를 받아들이는 과정을 세 단계로 설명했다.

첫 번째는 상실감. 익숙했던 것을 잃어버리는 슬픔이다. “예전 방식이 좋았는데.”

두 번째는 혼란. 옛 방식은 버렸는데 새 방식은 아직 익숙하지 않은 불안한 중간 지대다. “뭘 어떻게 해야 하는 거지?”

세 번째는 새로운 시작의 수용. 새 방식에 적응하고, 그 안에서 자기 자리를 찾는 것이다.

리더는 직원들이 첫 번째와 두 번째 단계에 있을 때, 무시하거나 서두르지 않아야 한다. “걱정 마, 다 잘 될 거야”라는 빈말 대신, “당신이 불안한 것을 안다. 당연한 거다. 우리 같이 넘어가자”라고 말해야 한다.

로마가 새로운 속주를 편입할 때, 현지인의 문화와 신을 존중해준 경우에는 안정적으로 다스릴 수 있었다. 하지만 현지인의 감정을 무시하고 로마의 방식을 일방적으로 강요한 경우(브리타니아의 부디카 반란이 대표적이다), 거센 저항에 부딪혔다. 사람의 감정을 무시하면 반드시 대가를 치른다.

I: 구조적인 발판을 세워라

감정적 이해만으로는 부족하다. 직원들이 실제로 새 방식을 익힐 수 있는 구체적인 지지대가 필요하다.

동영상 교육을 한두 번 보여주는 것으로는 안 된다. 직접 만져보고 실험할 수 있는 시간을 줘야 한다. 각 부서의 실제 업무에 맞는 사례를 보여줘야 한다. 그리고 현장에서 동료들이 신뢰하는 사람을 ‘AI 도우미(AI 챔피언)’로 지정해서, 옆에서 도와줄 수 있게 해야 한다.

새 방식이 자리를 잡으면, 성과를 평가하는 기준(KPI)과 보상 체계도 바뀌어야 한다. 옛날 방식으로 성과를 매기면서 새 방식으로 일하라고 하면, 아무도 움직이지 않는다. 보상이 가리키는 방향으로 사람은 움직인다.

F: 실수해도 괜찮다는 분위기를 만들어라

5장에서 깊이 다뤘던 심리적 안전감이 여기서도 핵심이다.

“챗GPT를 어떻게 쓰는지 물어보면 바보처럼 보일까 봐” 두려워서 묻지 못하는 순간, 혁신은 멈춘다. 리더가 먼저 “나도 모른다, 같이 배우자”라고 솔직하게 말해야 한다. “나도 어제 AI 프롬프트를 짜다가 완전히 망했다”라고 자기 실패를 공개해야 한다.

이 한마디가 만드는 효과는 측정 가능하다. 실패를 허용하는 팀에서 새 아이디어가 27% 더 나오고, 결정 속도가 35% 빨라지고, 핵심 인재의 이직률이 50% 줄어든다는 연구가 있다.

T: 저항을 적이 아닌 데이터로 활용하라

유럽의 한 대형 보험 회사 이야기를 하겠다. 이 회사가 보험 심사에 AI를 도입하려 하자, 경력 많은 심사역들이 강하게 반발했다. 기술이 싫어서가 아니었다. AI 시스템을 설계하는 과정에서 자기들의 전문 지식이 완전히 무시되었기 때문이었다. "20년 경력의 내 판단을 무시하고, 기계에게 맡기겠다고?"

리더십이 방법을 바꿨다. AI를 위에서 강제로 내리는 대신, 심사역들과 함께 AI 업무 과정을 설계하기 시작했다. AI가 심사역의 판단을 '대체'하는 게 아니라, 방대한 서류를 요약해서 심사역의 판단을 '돕는' 구조로 바꿨다. 심사역들의 전문성이 존중받자, 참여도가 폭발적으로 올라갔다. 이 부서의 성공이 회사 전체로 퍼지는 모범 사례가 되었다.

핵심 교훈이 있다. 직원의 저항은 적이 아니라, 시스템이 보내는 신호다. "이 부분이 잘못됐어요"라는 피드백이다. 이 신호를 무시하면 전쟁에서 정찰 보고를 무시하는 것과 같다. 카이사르는 갈리아 부족장들의 불만을 무시하지 않았다. 불만의 원인을 파악하고, 가능하면 해소하고, 불가능하면 우회했다. 저항을 힘으로 누르는 것은 최후의 수단이었다.

숫자로 증명하라: 겉치레 지표를 버리고 진짜 성과를 측정하라

조직 문화를 바꾸는 것과 함께, 성과를 어떻게 측정할지도 바뀌어야 한다.

많은 회사가 "AI 도구에 몇 명이 로그인했는가"를 센다. 이건 겉치레 숫자(허영 지표)다. 로그인만 하고 아무것도 안 하는 사람도 수에 포함되니까.

진짜 봐야 할 것은 두 층이다.

앞선 지표. AI 시스템이 건강하게 돌아가고 있는지를 본다. 사용자 수, AI가 자동으로 처리한 일의 양, AI가 엉뚱한 답을 내놓는 비율(환각 발생률), 직원들이 AI를 얼마나 자신 있게 쓰는지.

뒤따르는 지표. 이 AI 활용이 실제로 사업 성과로 이어졌는지를 본다. 고객 만족도가 올랐는가, 업무 처리 시간이 줄었는가, 실수가 줄었는가, 영업 이익이 늘었는가.

예를 들어, 개발자가 코드를 쓸 때 AI의 도움을 받는 것(앞선 지표)이, 제품 개발 기간 단축(뒤따르는 지표)으로 이어진다는 것을 증명해야 한다. 이 연결 고리를 보여주지 못하면, AI 프로그램은 예산을 잃고 사라진다.

6.2 AI를 안전하게 쓰면서 회사의 핵심 자산을 지키는 법

속도와 안전, 둘 다 잡아야 한다

로마의 도로는 빠르게 달릴 수 있도록 곧고 넓게 깔렸다. 하지만 도로만 있었던 것은 아니다. 도로의 요소요소에 초소가 있었고, 도적이 출몰하는 구간에는 순찰대가 돌았다. 빠르게 달릴 수 있는 도로와, 안전하게 달릴 수 있는 치안. 로마는 둘 다를 갖추었기에 제국 전역에서 물자와 사람이 자유롭게 이동할 수 있었다.

AI 시대에도 같은 원리가 적용된다. 빠르게 AI를 쓸 수 있으면서, 동시에 안전하게 통제할 수 있어야 한다.

2026년 현재, 기업용 AI 프로젝트의 42%가 실패하는데, 보안과 관리 체계가 처음부터 설계에 반영되지 않은 것이 큰 원인 중 하나다. 78%의 직원이 회사 몰래 개인 AI를 쓰는 '새도우 AI' 문제도 여전하다. 이 중 절반 가까이가 회사의 비밀 정보를 외부 AI에 입력하고 있다.

2025년 말, EU(유럽연합)와 미국 사이의 데이터 전송 협정이 유럽 법원에 의해 무효화되었다. 34개 이상의 나라가 데이터가 어디서 처리되는지를 규제하는 법률을 강화하고 있다. EU AI법(EU AI Act)은 2026년 8월부터 고위험 AI 시스템에 대한 투명성과 데이터 관리 요구사항이 본격 시행된다. 예전처럼 "그냥 클라우드에 올리면 되지"가 통하지 않는 세상이 되었다.

매번 사람이 확인하는 건 답이 아니다

AI를 안전하게 쓰려고, AI가 하는 모든 일에 사람이 승인 도장을 찍게 하는 회사가 많다. 겉보기에는 안전해 보인다. 하지만 이러면 AI를 쓰는 의미가 없어진다. 사람이 병목이 되어 전체 속도가 사람 속도로 묶인다.

2장과 3장에서 반복했던 이야기다. '절차적 안전장치' 대신 '구조적 안전장치'가 필요하다.

구조적 안전장치란, 시스템을 만들 때 처음부터 안전 규칙을 설계에 녹여넣는 것이다. 세 가지 방법이 있다.

첫째, 권한 상속. 회사에서 재무 데이터를 볼 수 없는 직원은, AI를 통해서도 그 데이터에 접근할 수 없다. 사람의 권한이 AI에게 그대로 넘어간다.

둘째, 정보 가리기. 직원이 외부 AI에 질문할 때, 중간에 보안 장치가 끼어들어서 고객 이름이나 정확한 금액 같은 비밀 정보를 자동으로 가린다. AI에게는 "X 고객에게 Y 금액을 제안하면"이 아니라, "가상 고객 A에게 가상 금액 B를 제안하면"으로 바뀌어 전달된다. AI의 두뇌는 빌리워, 비밀은 내보내지 않는다.

셋째, 확신도에 따른 분류. AI가 95% 이상 확신하는 건은 자동 처리하고, 확신이 낮은 5%의 애매한 건만 사람에게 넘긴다. 사람은 정말 사람의 판단이 필요한 어려운 경우에만 집중할 수 있다.

IBM이 제시한 AI 관리의 5가지 기둥

IBM은 자율적으로 움직이는 에이전트 AI를 안전하게 관리하기 위해 기업이 갖춰야 할 5가지 기둥을 내놓았다.

첫째, 정렬. AI가 회사의 가치와 목적에 맞게 행동하는지 확인하는 것이다. AI가 원래 목표에서 벗어나는 '표류 현상'을 실시간으로 감시한다.

둘째, 통제. AI가 접근할 수 있는 데이터와 도구의 범위를 엄격히 제한한다. AI가 폭주할 경우를 대비해 긴급 정지 버튼(킬 스위치)을 반드시 설계에 포함한다.

셋째, 추적. AI가 내린 모든 결정의 기록을 남긴다. 문제가 생겼을 때, "왜 이런 결정을 했지?"를 추적할 수 있어야 한다.

넷째, 보안. 외부 해커가 AI에게 잘못된 지시를 넣는 공격(프롬프트 인젝션)이나, 악의적인 데이터를 주입하는 공격으로부터 시스템을 보호한다. AI는 회사 방화벽과 분리된 안전한 구역(샌드박스)에서

실행되어야 한다.

다섯째, 법 준수. AI의 결정에 대해 누가 법적 책임을 지는지 명확히 한다. EU AI법 같은 최신 법규를 AI가 스스로 검토하는 장치를 넣는 것도 앞서가는 방법이다.

“당신의 질문이 가장 비싼 자산이다”

여기서 더 깊은 층위의 이야기를 하겠다.

엔비디아(NVIDIA)의 CEO 젠슨 황(Jensen Huang)은 2026년 1월 다보스 세계경제포럼에서 이렇게 말했다. “모든 나라는 자기만의 AI 인프라를 지어야 합니다. 자기 나라의 언어와 문화라는 천연자원을 활용해서 자기만의 AI를 만들고, 자기만의 국가 지능을 키워야 합니다.”

이 말은 나라에만 해당하는 게 아니다. 기업에도 똑같이 적용된다.

지난 10년간 IT 업계의 상식은 “모든 것을 클라우드에 올려라”였다. 자기 서버를 갖는 건 구식이라고 여겨졌다. 하지만 AI 시대에 이 상식이 흔들리고 있다.

왜 그런지 예를 들어보겠다. 글로벌 확장을 준비하는 제조업 사장이 외부 AI에 이런 질문을 한다고 하자. “다음 분기에 인도 시장에 진출하려는데, 2억 달러 예산으로 현지 공장을 세울 때 가장 큰 장벽은 뭔가? 경쟁사 A의 점유율을 뺏으려면 우리 신소재 기술을 어떻게 내세워야 할까?”

AI가 내놓는 대답은 좋을 것이다. 하지만 문제는 다른 곳에 있다. 이 질문 자체가 회사의 가장 비밀스러운 전략을 담고 있다. “인도에 2억 달러를 투자한다”, “경쟁사 A를 겨냥한다”, “신소재 기술이 무기다.” 이 정보가 외부 서버로 넘어간 것이다.

AI 시대에 대답은 누구나 얻을 수 있는 것이 되었다. 진짜 비싼 것은 질문이다. 질문은 곧 회사가 무엇을 고민하고 있고, 어디가 약한지를 보여주는 지도다. 경쟁사가 이 지도를 손에 넣으면, 코드를 도둑맞는 것보다 더 치명적이다.

젠슨 황은 2026년 5월 서비스나우(ServiceNow) 행사에서 한 발 더 나갔다. 에이전트 AI에 필요한 컴퓨팅 양이 기존 생성형 AI보다 최소 1,000%(10배) 이상 늘어날 것이라고 밝혔다. 그의 비전은 “100억 개의 디지털 AI 에이전트가 인간 직원과 함께 일하는 세상”이다. 이 규모의 AI가 회사 밖의 서버에서 돌아가면서 회사의 모든 질문과 데이터를 처리한다고 상상해보라. 위험은 기하급수적으로 커진다.

핵심은 우리 집에, 나머지는 빌려 쓴다

그래서 앞서가는 기업들은 ‘하이브리드 전략’을 택하고 있다.

회사의 모든 AI를 자기 서버에서 돌리는 건 비용이 너무 많이 든다. 반대로 모든 것을 외부 클라우드에 맡기면 비밀이 새어나갈 위험이 크다. 둘 다 극단이다. 답은 중간에 있다.

회사의 일을 세 등급으로 나눈다.

A등급. 회사의 생존과 직결되는 핵심 비밀, 독자적인 기술, 핵심 고객 데이터, 전략 알고리즘. 이걸 아무리 비용이 들어도 자기 서버(온프레미스)에서 처리한다. 금고 안에 넣어두는 것이다.

B등급. 민감하지만 비밀 정보를 가려낸 후에는 외부에서 처리해도 되는 일. 고객 이름과 금액을 가린 채로 클라우드 AI에 분석을 맡긴다.

C등급. 일반적인 업무. 이메일 번역, 마케팅 초안 작성, 회의록 요약 같은 것. 이건 외부 클라우드의 최신 AI를 마음껏 쓴다. 비밀이 들어있지 않으니까.

핵심은 자기 집에서 키우고, 주변 일은 남의 힘을 빌린다.

로마 제국의 방어 체계도 비슷했다. 로마 시의 방어는 직접 정예 군단이 맡았다. 변방의 보조적인 방어는 동맹군이나 용병에게 맡겼다. 핵심은 직접 지키고, 주변은 위임하는 것. 모든 곳을 정예 군단으로 채우면 비용이 감당이 안 되고, 모든 곳을 용병에게 맡기면 위험하다. 균형이 답이다.

데이터가 진짜 해자다

마지막으로 이 모든 전략의 궁극적 목적을 이야기하겠다.

AI 시대 초기에는 GPT-4나 클로드 같은 AI 모델 자체가 경쟁력인 것처럼 보였다. 하지만 시간이 지나면서 모델의 성능 차이가 줄어들고 있다. 누구나 비슷한 수준의 AI를 쓸 수 있게 되었다. 알고리즘은 평준화된다.

그러면 뭐가 남는가. 데이터다.

글로벌 자산운용사 블랙록(BlackRock)이 11조 달러 넘는 자산을 운용하면서 쌓아온 수십 년간의 시장 데이터. 대형 병원이 수백만 명의 환자를 치료하면서 쌓아온 임상 기록. 이런 데이터는 인터넷에서 긁어모을 수 없다. 복제할 수 없다. 이것이 진짜 해자(壕字)다. 해자란 성 주위를 둘러싼 물길, 적이 넘어올 수 없는 방어선이다.

BCG의 10:20:70 법칙을 2장에서 다뤘다. 알고리즘 10%, 기술 인프라 20%, 사람과 프로세스 70%. 여기에 하나를 더 보태야 한다. 이 모든 것 위에 데이터라는 지반이 있다. 아무리 좋은 건물(알고리즘)을 짓고, 좋은 사람(프로세스)을 모아도, 지반(데이터)이 약하면 건물은 기울어진다.

기업은 자기만의 데이터를 안전한 곳에 쌓고, 그 데이터로 AI를 자기 회사에 맞게 훈련시켜야 한다. 이것이 경쟁자가 돈으로 살 수 없는 유일한 경쟁력이다.

브레이크가 좋아야 빨리 달릴 수 있다

많은 사장이 AI 보안과 관리 체계를 “혁신의 속도를 늦추는 귀찮은 규제”로 여긴다.

정반대다.

F1 경주차가 시속 350킬로미터로 트랙을 질주할 수 있는 이유는 뭘까. 엔진이 좋아서? 물론이다. 하지만 진짜 이유는, 그 차에 가장 강력하고 신뢰할 수 있는 브레이크가 달려 있기 때문이다. 브레이크를 믿을 수 있으니까, 운전수는 직선 코스에서 가속 페달을 끝까지 밟을 수 있다. 브레이크가 없는 차는 빨리 달릴 수 없다. 무섭기 때문이다.

AI 거버넌스도 마찬가지다. 보안과 관리 체계가 단단한 회사의 직원들은, 실수나 정보 유출의 공포 없이 가장 대담하게 AI를 실험하고 활용할 수 있다. “이거 해도 되나?”를 매번 고민하지 않아도 되니까.

시스템이 이미 안전하게 설계되어 있으니까.

안전장치는 속도를 늦추는 것이 아니다. 속도를 가능하게 하는 것이다.

로마의 도로가 빠른 이동을 가능하게 한 것은 도로 자체의 넓이만이 아니었다. 도로를 안전하게 지키는 치안 체계, 일정한 간격으로 배치된 역참(驛站), 표준화된 도로 규격이 함께 작동했기에 사람과 물자가 빠르고 안전하게 이동할 수 있었다. 도로와 치안은 별개가 아니라 하나의 시스템이었다.

2026년, 파일럿의 연옥에서 빠져나와 AI를 회사 전체로 넓히려는 기업이 기억해야 할 것은 이것이다. 기술도 중요하고, 사람의 마음도 중요하고, 안전 체계도 중요하다. 이 셋은 따로 노는 것이 아니라, 하나의 시스템으로 엮여야 한다. 로마의 도로, 법, 군단이 하나의 제국을 이루었듯이.

KIMKJ.COM

마지막 글. 칼을 잡을 것인가, 전쟁을 설계할 것인가

이 책을 쓰는 동안, 세상은 내가 쓰는 속도보다 빨리 변했다.

1장을 쓸 때 소개한 AI 도구가, 6장을 쓸 때쯤 이미 구버전이 되어 있었다. 커서(Cursor)의 매출이 3개월 만에 두 배가 되고, 세일즈포스가 100개 넘는 AI 도구를 한꺼번에 쏟아내고, 소프트웨어 회사의 주가가 하루 만에 400조 원이 증발하는 세상이다. 이 속도 앞에서 어떤 책도 최신 상태를 유지할 수 없다.

그래서 나는 이 책에서 기술의 세부 사항보다, 변하지 않는 원리에 집중하려 했다.

로마인 이야기를 15권까지 읽고 나면, 독자의 머릿속에 남는 것은 개별 전투의 전술이 아니다. “왜 어떤 나라는 흥하고, 어떤 나라는 망하는가”라는 근본적인 질문에 대한 감각이다. 시오노 나나미는 로마가 강했던 이유를 군사력이 아닌 문화에서 찾았다. 패배한 적에게도 시민권을 주는 개방성, 실패를 부끄러워하지 않는 실용주의, 법 앞에 모두가 평등하다는 원칙.

이 책에서 반복해서 한 이야기도 결국 같은 것이다.

기술은 돈으로 살 수 있다. 좋은 AI 모델을 구매하고, 빠른 서버를 설치하는 것은 예산의 문제다. 하지만 기술을 제대로 쓸 수 있는 사람과 문화는 돈으로 살 수 없다. 실패를 허용하는 분위기, 리더가 먼저 뛰어드는 용기, 직원의 두려움에 솔직하게 답하는 정직함, 부서의 벽을 허무는 유연함. 이것들은 시간과 의지로만 만들어진다.

BCG의 10:20:70 법칙이 말하는 것도 이것이다. AI 알고리즘이 10%, 기술 장비가 20%, 사람과 일하는 방식이 70%. 승부의 70%는 기술 밖에 있다.

나는 이 책의 1장에서 체스 비유를 들었다. 조만간 모든 사람이 똑같이 강력한 AI라는 기물을 손에 쥐게 된다. 그때 승부를 가르는 것은 기물 자체가 아니라, 기보를 설계하는 능력이다. 6장을 마무리하는 지금도 같은 말을 되풀이한다. 결국 중요한 것은 기계가 아니라 사람이다.

카이사르의 갈리아 전쟁기를 읽으면, 그가 직접 칼을 휘두르는 장면은 거의 나오지 않는다. 대신, 어느 군단을 어디에 보낼지, 보급선을 어떻게 유지할지, 적의 내분을 어떻게 이용할지를 끊임없이 설계하는 모습이 나온다. 칼을 잡는 것은 병사의 일이었다. 전쟁을 설계하는 것이 카이사르의 일이었다.

AI 시대에 묻는다. 당신은 칼을 잡을 것인가, 전쟁을 설계할 것인가.

실행은 점점 더 AI의 몫이 되어간다. 데이터를 입력하고, 보고서를 쓰고, 코드를 짜는 일. 이런 ‘칼 휘두르기’는 기계가 더 빠르고 정확하게 해낸다. 사람에게 남는 것은 ‘왜 이 전쟁을 하는가’를 정하는 일이다. 목적을 부여하고, 방향을 잡고, 기계가 놓치는 인간적인 부분을 채우고, 결과에 책임을 지는 일.

이 책이 독자에게 드리고 싶은 것은 기술 매뉴얼이 아니다. 변화 앞에서 자기 자리를 찾는 감각이다. 로마는 하루아침에 이루어지지 않았다. 당신의 변화도 그럴 것이다. 하지만 첫 걸음을 떼는 것은 오늘 할 수 있는 일이다.

이 책을 읽고, 내일 아침 AI를 한번 직접 써보시길 바란다. 서투르더라도 괜찮다. 리더든 신입이든, 첫 걸음은 누구에게나 서투르다. 카이사르도 처음에는 무명의 정치인이었다. 중요한 것은 시작하는 것이다.

KIMKJ.COM

이 책을 잘 읽으셨으면 그리고 새로운 가치있는 지식을 얻으셨다고 판단되시면
농협 302-1096-0948-81 (예금주 김경진) 에 자발적 후원 부탁드립니다.