

신뢰의 설계자들

두 개의 AI법

유럽과 한국은 같은 기술을

어떻게 다르게 붙잡았는가

신뢰의 설계자들

두 개의 AI법 — 유럽과 한국은 같은 기술을 어떻게 다르게 붙잡았는가

2026년 여름, 같은 기술을 향해 두 개의 법이 거의 동시에 작동을 시작했다.
이 책은 그 두 그물이 어디서 겹치고 어디서 어긋나는지를 조문 단위로
따라간다.

신뢰의 설계자들

부제	두 개의 시법 — 유럽과 한국은 같은 기술을 어떻게 다르게 붙잡았는가
원고 기준일	2026-08-07
통합본 생성일	2026-08-12
분량	본문 95,324자 (프롤로그·0부~4부·에필로그·후기·부록)
구성	프롤로그 1 · 준비 장 3 · 본문 15장 · 에필로그 · 후기 · 부록 4

내부 검토본. 이 문서는 원고 통합·조판 결과물이며 외부 발행본이 아니다. 수록된 정책 서술을 외부에 배포하거나 공표하기 전에는 소속기관의 사전 검토 대상인지 확인이 필요하다.

본문이 다루는 두 법은 지금도 움직이고 있다. **[검토 필요]** 표시가 붙은 문장은 1차 원문으로 확인되지 않은 대목이므로 인용하지 말고 부록 D의 원문을 직접 확인할 것.

이 책을 읽는 방법

처음부터 읽는 독자에게. 0부를 건너뛰지 않기를 권한다. 두 AI법의 조문은 그 자체로는 낯선 외국어에 가깝다. 0부 세 편은 그 외국어의 문법에 해당한다. 세 편을 합쳐 한 시간이 걸리지 않는다.

필요한 것만 찾는 독자에게. 각 장은 첫머리에 **"이 장을 한 문장으로"** 상자를 두었다. 목차와 그 상자만 훑어도 어디를 펴야 할지 알 수 있다. 장 끝에는 **"실무자를 위한 체크박스"**가 구분선 아래 따로 붙어 있다. 다만 필요하다면 거기서부터 거꾸로 읽어도 된다.

실무 담당자에게. 4부가 당신의 자리다. 다만 12장의 분류 결정 트리는 3장 (고영향 vs 고위험)을, 15장의 조달 논의는 7장(검·인증과 우선구매)을 전제로 한다. 두 장만 먼저 읽고 4부로 건너뛰는 경로를 권한다.

확정되지 않은 것에 관하여. 이 책이 다루는 두 법은 지금도 움직이고 있다. 한국은 상당수 세부 기준이 하위 고시를 기다리는 중이고, 유럽은 시행 일정 자체가 다시 조정되는 중이다. 그래서 본문은 **"법령으로 확정된 것"**과 **"아직 정해지지 않은 것"**을 문장 단위로 구분해 적었다. 확인되지 않은 대목에는 [검토 필요]를 붙였다. 그 표시가 붙은 문장은 인용하지 말고, 부록 D의 원문을 직접 확인하기 바란다.

차 례

프롤로그 — 신뢰는 자연발생하지 않는다	1
-----------------------------	---

우리는 얼굴을 모르는 사람이 만든 물건에 매일 목숨을 맡긴다. 이제 맡기는 것은 물건이 아니라 판단이다.

0부 — 두 법을 읽기 전에

유럽이 40년에 걸쳐 그린 설계도를 모르면, AI법이 왜 이런 모양인지 보이지 않는다.

0-1 유럽은 왜 AI를 ‘제품’으로 다루는가	6
---------------------------------	---

법전에 모든 사양을 적어 넣던 방식은 어떻게 무너졌는가.

0-2 CE 마크 뒤의 40년 — 모듈·검증기관·시장감시	11
---------------------------------------	----

누가 확인하고, 팔린 뒤에는 누가 지켜보는가.

0-3 같은 법을 읽고 왜 다른 결론이 나오는가	17
----------------------------------	----

법령해석의 네 가지 방법. 이 책이 두 법을 읽는 방식에 대한 선언.

1부 — 설계도 대조: 두 법의 뼈대

같은 기술을 향한 두 개의 그물. 그물코가 왜 이렇게 다른가.

1장 두 법은 다른 질문에서 태어났다	22
----------------------------	----

진흥과 윤리를 함께 담은 법, 위험만을 기준으로 삼은 법. 제정이유서로 읽는 설계 철학.

2장 AI Act는 새 법인가, 기존 법의 확장인가	30
------------------------------------	----

40년 된 건물 위에 올린 새 층. 무엇이 이어졌고 무엇이 끊어졌는가.

3장 무엇을 위험이라 부르는가 — 고영향 vs 고위험	37
-------------------------------------	----

열 개 영역의 단일 목록과, 두 갈래로 뺀 나무.

4장 아예 금지된 것 — 10개와 0개	45
-----------------------------	----

금지 목록을 가진 법과, 그 목록 자체가 없는 법.

5장 살아 움직이는 제품 — 업데이트와 재인증	51
공장을 떠난 뒤에도 변하는 것을, 한 시점에 고정하는 인증으로 붙잡을 수 있는가.	

2부 — 의무의 실제: 무엇을 해야 하는가

돈과 시간이 실제로 드는 지점들.

6장 AI가 만든 것임을 밝히는 두 방식	59
사람에게 말로 알리라는 법과, 기계가 읽을 수 있게 새기라는 법.	

7장 자율의 얼굴을 한 규제 — 검·인증과 우선구매	67
“민간 자율”이라 불리는 조항이 조달이라는 지렛대를 만나면 무슨 일이 생기는가.	

8장 가장 큰 모델의 문턱 — 10^{26} 과 10^{25} 사이	77
지수 하나 차이가 만드는 열 배의 간극, 그리고 문턱 너머의 의무.	

9장 어지면 어떻게 되는가 — 3천만원 vs 매출 7%	85
제재의 크기에서 가장 노골적으로 드러나는 철학의 차이.	

3부 — 사고 이후: 책임과 감시

잘못됐을 때 누가 물어주고, 평소에는 누가 지켜보는가.

10장 AI가 잘못 판단했을 때 누가 책임지는가	93
피해자가 입증할 수 없는 것을, 법은 누구에게 지웠는가.	

11장 누가 감시하는가 — AI Office와 과기정통부	100
중앙 집행기구를 새로 만든 쪽과, 만들지 않은 쪽.	

4부 — 기업의 눈: 규제를 사업으로 번역하기

이 부는 실무자를 위한 것이다. 순서는 그 자리에서 실제로 일이 흘러가는 순서를 따랐다.

12장 우리 서비스는 어느 칸에 들어가는가	108
분류 결정 트리 — 두 법 중 어디에, 어느 등급으로 걸리는가.	

13장 두 법을 동시에 맞추는 최소 공통 작업	114
---------------------------------	-----

중복을 줄이는 순서. 무엇을 먼저 하면 나머지가 따라오는가.

14장 정책 환경을 읽는 법 — 어디를 보고 무엇을 기다리나	120
---	-----

모니터링 지점의 목록과, 아직 나오지 않은 것들의 대기 목록.

15장 규제를 사업 기회로 — 조달·표준·인증 시장	126
------------------------------------	-----

비용이 아니라 진입장벽으로 읽을 때 보이는 것들.

에필로그 — 어느 쪽이 더 똑똑한 규제인가	130
-------------------------------	-----

답을 고르는 대신, 40년 후에도 작동할 규제인가를 묻는다.

후기 — 규제 지식베이스는 어떻게 만드는가	135
-------------------------------	-----

이 책을 만든 방법. 원문 수집에서 분석 문서까지, 재현 가능한 절차로.

부록

조문 대응, 시행 일정, 용어, 원문 출처.

A. 한국 인공지능기본법 ↔ EU AI Act 조문 대응표	140
--	-----

B. 시행 일정 캘린더 2026~2030	143
------------------------------	-----

C. 용어 사전	148
----------------	-----

D. 원문·출처 목록	152
-------------------	-----

프롤로그 — 신뢰는 자연발생하지 않는다

우리는 매일 아침, 얼굴을 모르는 사람들과 수십 번 악수한다.

커피 머신의 스위치를 누를 때 우리는 손끝으로 전류가 흘러들지 않으리라 믿는다. 아이 손에 나무 블록을 쥐여줄 때 모서리가 살을 베거나 철이 벗겨져 그 입으로 들어가지 않으리라 믿는다. 지하철에 오를 때 이 수십 톤짜리 쇳덩이의 제동장치가 오늘 아침에도 제 일을 하리라 믿는다.

우리는 그것들을 만든 사람의 얼굴을 모른다. 공장이 어느 나라에 있는지, 어떤 설계도로 만들었는지, 마지막으로 누가 검사했는지 모른다.

그런데도 목숨을 맡긴다.

이 믿음은 어디서 오는가

많은 사람이 이 믿음이 기업의 양심이나 시장의 평판에서 온다고 여긴다. 혹은 국가라는 파수꾼이 뒤를 지켜주기 때문이라고 생각한다.

둘 다 절반만 맞다.

우리가 매일 아무 생각 없이 누리는 이 신뢰는 저절로 생겨난 것이 아니다. 그것은 설계된 것이다. 누군가 앉아서 서로 얼굴을 볼 일이 없는 수억 명의 인간이 어떻게 하면 서로를 해치지 않고 거래할 수 있을지를 계산해서 그려낸 도면이 있다.

그 도면의 이름은 대부분의 사람에게 낯설다. 그러나 도면이 남긴 흔적은 우리 집 안 어디에나 있다. 충전기 뒷면, 세탁기 옆구리, 장난감 상자 구석에 찍힌 두 글자.

CE.

금속 딱지나 스티커에 찍힌 이 두 글자 뒤에는 40년에 걸친 하나의 질문이 압축돼 있다. 서로 다른 언어를 쓰고 서로 다른 이해관계를 가진 스물일곱 나라의 기업과 소비자가, 한 번도 만나지 않고 어떻게 서로를 믿을 것인가.

그런데 이제 말기는 것이 달라졌다

지금까지 우리가 얼굴 모르는 상대에게 맡긴 것은 물건이었다.

물건에는 몸이 있다. 무게가 있고 재질이 있고 만져지는 표면이 있다. 몸이 있으면 검사할 수 있다. 전압을 재고, 화학물질을 뽑아 분석하고, 떨어뜨려 보고, 부러뜨려 본다. 결정적으로 몸이 있는 물건은 공장을 떠난 뒤로 변하지 않는다. 오늘 검사에 합격한 그 커피 머신은 3년 뒤에도 같은 커피 머신이다.

이제 우리가 맡기는 것은 물건이 아니라 **판단**이다.

수천 통의 이력서 중 내 것이 사람 눈에 닿을지를 결정하는 판단. 대출 신청서를 승인과 거절로 가르는 판단. CT 영상에서 그림자 하나를 종양이라 부를지 아무것도 아니라 부를지를 정하는 판단. 한밤중 상담창 너머에서 내 질문에 답하는 것이 사람인지 아닌지 알 수 없는 상대가 내리는 판단.

이 판단에는 몸이 없다. 떨어뜨려 볼 표면도, 뽑아낼 화학물질도 없다. 무엇보다 이것은 공장을 떠난 뒤에도 계속 변한다. 어젯밤 업데이트로 오늘의 판단은 어제의 판단이 아니다.

40년 동안 몸을 가진 물건을 전제로 그려온 신뢰의 설계도가 처음으로 흔들리는 지점이 여기서다.

2026년 여름, 두 개의 그물

그래서 두 개의 법이 만들어졌다.

2026년 7월 21일, 한국에서 「인공지능 발전과 신뢰 기반 조성 등에 관한 기본법」이 전면 시행됐다.^[1] 그로부터 열이틀 뒤인 8월 2일, 유럽연합에서 AI Act의 투명성 의무가 적용되기 시작했다.^[2]

같은 계절, 같은 기술을 향해 두 개의 그물이 거의 동시에 던져진 것이다.

그런데 두 그물은 닻지 않았다.

한쪽은 열 가지 인공지능을 아예 금지 목록에 못 박았고, 다른 쪽은 그런 목록을 만들지 않았다. 한쪽은 여기면 전 세계 매출의 7퍼센트를 물리고, 다른 쪽은 3천만 원을 물린다. 한쪽은 법의 이름에 "발전"과 "신뢰"를 나란히 넣었고, 다른 쪽은 이름에 아무 약속도 넣지 않았다.

같은 기술을 붙잡으려는 그물의 코가 이렇게까지 다를 수 있는가.

이 책이 묻는 것

이 책은 어느 쪽이 옳은지를 가리려는 것이 아니다.

법조문 두 벌을 나란히 놓고 다른 곳에 밑줄을 긋는 일은 어렵지 않다. 어려운 것은 그다음이다.

왜 다르게 설계되었는가. 어떤 질문에 답하려다 이런 모양이 되었는가. 그리고 그 모양은 앞으로 40년을 버틸 수 있는가.

그 답을 찾으려면 2026년 여름에서 출발해서는 안 된다. 1985년, 유럽이 법전에 제품의 모든 사양을 적어 넣는 일을 포기한 그 순간까지 거슬러 올라가야 한다. 그때 유럽이 발명한 문법이 오늘 AI Act의 뼈대이기 때문이다.

그래서 우리는 준비운동부터 시작한다. CE 두 글자 뒤에 무엇이 있는지를 먼저 본 다음, 두 개의 AI법을 조문 단위로 나란히 펼칠 것이다.

이 이야기의 끝에서 우리가 마주할 질문은 아마 이것일 테다.

우리는 지금까지 얼굴을 모르는 사람이 만든 물건을 믿는 방법을 배워왔다. 이제 우리가 얼굴을 모른 채 신뢰해야 하는 것은 물건이 아니라 판단이다.

그것을 믿는 방법은, 아직 아무도 다 설계하지 못했다.

주(註)

1. 법률 제21311호. 2026년 1월 22일 공포된 시행령(대통령령 제36053호)과 7월 21일 시행된 개정 시행령(제36506호)이 함께 작동한다. 다만 우선구매 확인·연산량 산정방식 등 상당수 세부 기준은 하위 고시로 미뤄져 있어, 이 책은 "법령으로 확정된 것"과 "고시를 기다리는 것"을 구분해 서술한다.
2. Regulation (EU) 2024/1689. 전체 113개 조문과 13개 부속서가 한꺼번에 발효한 것이 아니라 Article 113에 따라 단계적으로 적용된다. 금지 AI는 2025년 2월, 범용 AI 모델(GPAI) 의무는 2025년 8월, 투명성 의무를 포함한 대부분의 조항은 2026년 8월이다. 이 일정은 이후 개정 논의로 다시 조정될 수 있어 [검토 필요] 표시를 붙인다.

PART 0

0부 — 두 법을 읽기 전에

유럽이 40년에 걸쳐 그린 설계도를 모르면, AI법이 왜 이런 모양인지 보이지 않는다.

0-1. 유럽은 왜 AI를 '제품'으로 다루는가

이 장을 한 문장으로

유럽은 1985년에 "어떻게 만들어야 하는가"를 법에서 지워버렸고, 그 빈자리를 메우려고 40년 동안 지은 구조물이 오늘 AI를 붙잡고 있다.

1985년 이전의 유럽을 상상해 보자.

독일의 계량기 제조사가 저울을 만든다. 독일이 요구하는 안전 기준을 모두 충족했다. 그런데 프랑스 국경에서 막힌다. 프랑스의 기준이 다르다. 이탈리아도, 스페인도 제각각이다.

당시는 더 심했다. 계량기 같은 정밀기기 분야에서는 정부 공무원이 직접 공장에 와서 제품 하나하나를 검사하고 도장을 찍었다.^[1] 기업이 아무리 잘 만들어도 정부가 확인하기 전까지는 팔 수 없었다.

국가가 심판이자 검사관이었다.

법전에 모든 것을 적어 넣던 시대

이것을 Old Approach(구점근방식)라고 부른다. 각국 정부가 시장을 신뢰하지 못해 모든 기술 규격을 법령에 직접 박아 넣고 스스로 검증하던 방식이다.

조선의 경국대전이 그랬다. 나라 운영 원리부터 관리의 복장, 제례 절차까지 법전 한 권에 다 담았다. 세상이 바뀌어도 법전은 그대로였다.

유럽의 Old Approach도 같은 운명이었다.

기술이 바뀌면 법이 뒤쳐졌다. 새로운 소재가 나왔다. 법에 없었다. 새로운 구조가 나왔다. 법에 없었다. 법을 고치는 데 몇 년이 걸렸고 그사이 기업은 멈춰 서 있었다.

스물일곱 나라가 하나의 시장을 만들려면 이 방식으로는 불가능했다.

1985년, 질문을 바꾸다

유럽은 답을 찾는 대신 질문을 바꿨다.

"어떻게 만들어야 하는가"를 법에 쓰지 말고, "무엇을 달성해야 하는가"만 쓰면 어떨까.

이것이 **New Approach**(신접근방식)다.

법령에는 제품이 달성해야 할 **필수 요구사항**(Essential Requirements)만 남긴다. "기계는 사용자를 다치게 해서는 안 된다." 두께 몇 밀리미터, 재질 무엇, 이런 내용은 없다. 목표만 있다.

방법은 기업의 몫이다.

여기에 절묘한 장치가 하나 붙는다. 유럽 표준화 기구가 "이렇게 만들면 법적 요건을 충족한다"는 구체적인 방법을 **조화표준**(Harmonised Standards)으로 제시하고, 그 표준을 따르면 국가가 일단 안전하다고 인정해 준다. 이것을 **준수 추정**(Presumption of Conformity)이라 한다.

[2]

표준을 따르지 않아도 된다. 다른 방식으로 안전을 입증하면 된다. 다만 그 경우에는 입증 책임이 기업에게 온다.

자유를 줬다. 책임도 함께 줬다.

남은 두 개의 질문

목표는 정해졌다. 그런데 곧바로 다음 질문이 따라왔다.

그 목표를 달성했다는 것을 어떻게 증명하는가.

1989년, 유럽은 **모듈(Module)** 시스템으로 답한다. 제품의 위험도와 기업의 규모에 따라 고를 수 있는 증명 방법의 메뉴판을 만든 것이다. 기업이 스스로 선언하는 가장 가벼운 Module A부터, 외부 기관이 제품 하나하나를 검사하는 Module G, 기업의 품질 시스템 전체를 인증받는 Module H까지 여덟 갈래가 있다.

원리는 단순하다. 위험이 낮으면 기업 스스로 한다. 위험이 높으면 제3자가 개입한다.

메뉴판이라고 했지만 손님이 아무것이나 고를 수는 없다. **어떤 모듈이 메뉴에 오르는지는 입법자가 정한다.** 제품군별 법령이 "이 제품은 이 모듈들 중에서 고르라"고 선택지를 미리 좁혀 둔다.^[3]

그리고 또 하나의 질문이 남았다.

그 증명을 수행하는 기관은 믿을 수 있는가. 시장에 나온 뒤에는 누가 감시하는가.

2008년, 유럽은 세 개의 규정으로 답한다. 검증기관의 자격을 심사하는 인정 체계, 시장에 나온 제품을 사후에 감시하는 체계, 그리고 앞으로 만들어질 모든 법령이 같은 용어와 같은 모듈 체계를 쓰도록 강제하는 공통 문법. 이 셋을 묶어 **NLF(New Legislative Framework, 신법률체계)**라 부른다.

세 이름이 헛갈리는 이유는 사실 하나의 건물이기 때문이다.

1985년에 무엇을 정했고, 1989년에 어떻게를 정했고, 2008년에 누가를 정했다. 시간 순서로 답이 쌓였을 뿐 별개의 체계가 아니다.

[필자 분석] 이 설계의 진짜 의도

NLF를 "유연한 규제"라고만 소개하는 글이 많다. 틀린 말은 아니지만 절반의 설명이다.

이 설계의 더 깊은 목적은 **규제의 내구성**에 있다.

법을 고치는 데는 몇 년이 걸린다. 기술은 그보다 훨씬 빨리 바뀐다. 법령에 기술 사양을 박아 두는 방식은 규제가 언제나 기술보다 한 발 뒤처지는 구조를 만든다. NLF는 법이 기술을 따라가게 하지 않았다. 대신 표준이 기술을 따라가게 했다. 법은 목표만 지키면 된다.

이 설계가 40년 뒤 AI에 그대로 적용된다.

EU AI Act가 "AI 시스템은 안전해야 하고 투명해야 한다"고 쓸 뿐 "모델의 파라미터는 몇 억 개 이하여야 한다"고 쓰지 않는 이유가 여기에 있다. 조문을 읽고 "너무 추상적이다"라고 느낀다면, 그것은 결함이 아니라 설계다.

그래서 이 장의 제목에 답하자면 이렇다. 유럽이 AI를 제품으로 여겨서 제품법의 문법으로 다루는 것이 아니다. **거꾸로다**. 유럽에는 새로운 기술이 나올 때마다 꺼내 쓰도록 40년에 걸쳐 다듬은 문법이 이미 있었고, AI가 그 문법으로 붙잡을 수 있는 첫 비물질이었을 뿐이다.

문법이 먼저 있었고, 대상이 나중에 왔다.

그 문법이 실제로 어떻게 돌아가는지 — 누가 확인하고, 팔린 뒤에는 누가 지켜보는지 — 를 다음 장에서 본다.

주(註)

1. Blue Guide 2022, Section 1.1.1. 법정계량 분야에서 Old Approach 하에 공공기관이 직접 적합성 증명서를 발급한 사례를 명시한다. Blue Guide는 법령이 아니라 EU 집행위원회가 관보에 공고 형식으로 게재한 공식 해설서이며 구속력이 없다.
2. Blue Guide 2022, Section 1.1.3.
3. Blue Guide 2022, Section 5.1. "모듈의 복잡성은 제품의 리스크에 비례해야 한다"는 원칙과, 공익 보호에 필요한 수준을 유지하는 선에서 제조사에게 가능한 한 넓은 선택지를 주어야 한다는 원칙을 함께 제시한다. 법적 근거는 Decision 768/2008/EC와 이를 인용한 부문별 조화법령이다.

0-2. CE 마크 뒤의 40년 — 모듈·검증기관·시장감시

이 장을 한 문장으로

검증은 민간 기관이 하고, 그 기관은 다시 검증받고, 팔린 뒤에는 별도의 감시망이 붙는다. 그 세 겹을 조절하는 눈금이 리스크다.

스포츠에는 두 종류의 전문가가 있다.

코치는 선수 편이다. 약점을 찾아주고 전략을 함께 세운다. 심판은 어느 편도 아니다. 규칙을 적용하고 위반을 판정한다. 심판이 코치처럼 굴면 경기가 무너진다.

EU 인증 체계에서 검증을 맡는 **검증기관**(Notified Body, NB)은 심판이다.

심판의 자격

NB는 특정 EU 법령 아래 적합성평가를 수행하도록 회원국 정부가 지정하고 EU 집행위원회에 통보한 제3자 기관이다. 세 단어가 중요하다.

제3자. 제조사도 정부도 아닌 독립 기관이다. 대부분 민간기업이다.

공식 지정. 기술력이 아무리 뛰어나도 정부의 심사를 거쳐 NANDO(EU 공식 검증기관 데이터베이스)에 등재되기 전에는 NB가 아니다.

특정 법령 아래. 범용 자격이 아니다. 기계류 인증을 할 수 있는 기관이 의료기기 인증은 못 할 수 있다.

독립성 요건은 엄격하다. NB는 "이렇게 하면 통과한다"는 방법을 설계해 주어서는 안 된다. 기술 문서가 요건을 충족하는지 판정하는 것이 역할이고, 그 문서를 어떻게 채울지는 제조사의 책임이다.

CE 마크 옆에 붙는 네 자리 숫자가 그 NB의 고유 번호다. 그리고 그 옆에 찍히는 것은 NB의 보증이 아니라 **제조사 자신의 적합성 선언**이다. 검증을 받았든 안 받았든 최종 책임은 만든 쪽에 있다.

심판도 속는다

2010년 3월, 프랑스 의약품 안전청이 발표한 사실이 유럽 의료기기 규제를 뒤집었다.

프랑스 회사 PIP가 유방 보형물에 의료용 실리콘 대신 **산업용 실리콘**을 몰래 써 왔다는 것이다. 검증기관 TÜV Rheinland가 2000년에 CE 마크를 내준 뒤 약 10년, 65개국 40만 명에게 팔린 제품이었다.^[1]

TÜV는 10년 동안 PIP 공장을 여덟 차례 감사했다. 그러나 모든 방문이 사전에 통보된 방문이었고, 원자재나 사업 기록을 직접 검사하지는 않았다.

TÜV의 잘못이었을까. EU 사법재판소는 2017년에 판단했다. **NB에게는 예고 없는 방문이나 제품 직접 검사에 대한 일반적 의무가 없다.**^[2] TÜV는 당시 법령이 요구하는 것을 다 했다. PIP는 감사관이 오기 전에 재료를 숨기고 서류를 갖춰 두었다. 형사재판에서 TÜV는 사기의 피해자로 인정받았다.

문제는 심판이 아니라 규칙이었다.

EU는 2017년 의료기기 규정을 전면 개정해 예고 없는 방문과 광범위한 문서 접근권을 NB 의무로 추가했다. 규제가 사고를 따라 고쳐진 것이다.

여기서 신뢰의 사슬 구조가 드러난다. 제조사를 NB가 검증하고, NB를 인정기구(Accreditation Body)가 검증하고, 인정기구를 정부가 감독한다. 어느 고리도 혼자 서 있지 않다. 그런데도 고의적 사기 앞에서는 사슬 전체가 뚫릴 수 있다.

인증은 입장권이지 보증서가 아니다

모듈도 NB도 제품이 시장에 나오기 전의 이야기다. 실제 위험은 나온 뒤에 드러난다.

소비자가 예상하지 못한 방식으로 쓴다. 시간이 지나며 재료가 삭는다. 소프트웨어가 업데이트되며 동작이 바뀐다. 시험실에서 안 나오던 결함이 실제 환경에서 나온다.

그래서 2019년, EU는 시장감시 체계를 전면 강화한다.^[3] 바뀐 것 중 셋이 중요하다.

1. 온라인 판매에 대한 직접 집행력 — EU 안에 사업장이 없는 제3국 제조사는 EU 내 책임자를 반드시 지정해야 함. 한국 기업도 여기에 해당함.
2. 회원국 간 즉시 공조 — 한 나라가 발견한 위험이 나머지 전체로 전파되는 절차를 법제화함.
3. 생애주기 전반 감시 — "처음 출시될 때"가 아니라 "유통되는 내내"가 감시 대상임. 창고에 쌓인 것도, 이미 소비자 손에 있는 것도 포함함.

이 체계를 실시간으로 돌리는 것이 **Safety Gate**(옛 이름 RAPEX)다. 독일 당국이 어느 완구에서 기준치를 넘는 납을 발견해 집행위원회에 통보하면, 24시간 안에 전 회원국에 경보가 나간다. 정보는 역외로도 공유된다. 한국 국가기술표준원도 이 경보를 받아 국내 조치 여부를 검토한다.

무게중심이 사전에서 사후로 옮겨 간 이유는 셋이다. 온라인 직구의 폭증, 제품의 소프트웨어화, 공급망의 복잡화. 출시 전에 모든 것을 잡아내는 일이 불가능해졌다면, 출시 후에 빨리 찾아내는 체계가 그 공백을 메워야 한다.

모든 것을 조절하는 눈금

여기까지 오면 하나의 변수가 계속 등장했음을 알아차리게 된다.

모듈의 무게도, NB의 개입 여부도, 경보의 우선순위도 전부 같은 것에 비례한다. **리스크**다.

그런데 규제에서 "위험"과 "리스크"는 다른 말이다.

칼날은 **위험원**(Hazard)이다. 피해를 일으킬 수 있는 잠재적 원천이며, 칼날 자체는 아직 아무도 다치게 하지 않았다. 손가락이 베이면 그것이 **위해**(Harm)다. 그리고 **리스크**는 이렇게 정의된다.

리스크 = 위험이 발생할 확률과 그 위험의 심각도의 조합^[4]

같은 칼이라도 요리사가 장갑을 끼고 쓰면 리스크가 낮고, 아이가 맨손으로 장난치면 높다. 위험원은 같은데 확률과 심각도의 조합이 다르다.

이 구별이 왜 결정적인가. 규제가 "위험한 것은 금지한다"고 쓰면 칼 자체를 금지해야 한다. 그러나 "리스크가 허용할 수 없는 수준이면 조치한다"고 쓰면, 칼은 허용하되 어린이용 제품에는 칼날 커버를 요구할 수 있다.

같은 표준은 "안전"도 정의한다. **허용할 수 없는 리스크가 남아 있지 않은 상태**다.^[5] 리스크가 0인 제품은 없다는 뜻이고, 그 허용선이 기술 발전에 따라 계속 올라간다는 뜻이기도 하다. 20년 전에는 에어백 없는 자동차도 안전했다. 지금은 아니다.

[필자 분석] 측정이 없으면 비레도 없다

40년의 설계를 한 단어로 줄이면 비레성이다. 위험이 크면 검증이 무겁고, 작으면 가볍다.

비레성이 없으면 두 극단이 생긴다. 모든 제품에 최고 수준의 검증을 요구하면 기업이 비용에 짓눌려 아무것도 만들지 못한다. 모든 제품에 자기 선언만 허용하면 PIP 같은 일이 반복된다.

그런데 비레성이 작동하려면 전제가 하나 있다. 리스크를 썰 수 있어야 한다. "위험하다/안전하다"의 이분법으로는 비레를 구현할 수 없다. 어느 정도 위험한지를 재야 그에 비레하는 규제를 설계할 수 있다.

"확률 × 심각도"라는 공식이 학술적 정의에 그치지 않는 이유가 여기 있다. 이 공식이 없으면 비레성이 무너지고, 비레성이 무너지면 40년 설계 전체가 무너진다.

이 대목을 기억해 두기 바란다.

한국 인공지능기본법의 영향평가 조항은 인공지능이 기본권에 미치는 부정적 영향을 식별하고 분석하라고 쓴다. 확률과 심각도를 계산하라고는 쓰지 않는다.^[6] 의사가 "아파 보인다"고 적는 것과 "혈압 180/110, 즉시 투약"이라고 적는 것의 차이다. 앞은 판단이고 뒤는 측정이다.

두 법이 같은 단어를 쓰면서 다른 일을 하고 있다는 첫 신호가 여기서 나온다. 3장에서 이 이야기를 다시 꺼낸다.

지금까지 우리는 유럽이 지은 구조물을 봤다. 그런데 이 구조물을 읽는 사람마다 결론이 갈린다면 어떻게 해야 하는가. 그것이 다음 장의 문제다.

주(註)

1. PIP 스캔들 경우는 EU 의료기기 규제 개정의 직접 계기가 되었다. 관련 판결과 보도는 blog Series A의 A-3에 상세 출처를 정리해 두었다.
2. EU 사법재판소 2017년 2월 16일 판결.
3. Regulation (EU) 2019/1020 (제품의 시장감시 및 적합성에 관한 규정), OJ L 169, 25.6.2019.
4. ISO/IEC Guide 51:2014, Clause 3.9. "Risk: combination of the probability of occurrence of harm and the severity of that harm." 이 공식은 EU AI Act Article 9에 직접 채용되었으며, 집행위원회 디지털총국과 공동연구센터가 2024년 5월 30일 웨비나에서 이를 공식 확인했다.
5. ISO/IEC Guide 51:2014, Clause 3.14. "Safety: freedom from risk which is not tolerable." 같은 표준 3.15 Note 1은 acceptable risk와 tolerable risk를 동의어로 규정한다. 허용 수준을 기술·지식 발전에 따라 재검토해야 한다는 요구는 Clause 6.2.2에 있다.
6. 인공지능 발전과 신뢰 기반 조성 등에 관한 기본법(법률 제21311호) 제35조(고영향 인공지능 영향평가). 이 조항은 "노력하여야 한다"는 노력 의무로 규정돼 있다. 2026-08-04 법제처 국가법령정보 목차로 대조 확인했다.

0-3. 같은 법을 읽고 왜 다른 결론이 나오는가

이 장을 한 문장으로

법의 문장은 스스로 완결되지 않는다. 두 법을 비교하려는 사람이 가장 먼저 챙겨야 할 것은 조문이 아니라 조문을 읽는 방법이다.

공원 입구에 안내판이 하나 서 있다. **탈것 출입 금지.**

자동차는 안 된다. 그건 분명하다. 그럼 자전거는. 유모차는. 전동휠체어를 탄 사람은 공원에 들어갈 수 없는가.

법철학자 하트가 1961년에 던진 이 질문은 60년이 지나도 유효하다.^[1] 언어의 그물은 성글고, 세상은 그 그물코 사이로 끊임없이 빠져나간다. 법이 언어로 짜여 있는 한 법을 읽는 일은 언제나 해석하는 일이다.

대법원이 정해 둔 읽기의 순서

한국 대법원은 법을 읽는 순서를 판례로 정해 두었다. 법해석의 목표는 "법적 안정성을 저해하지 않는 범위 내에서 구체적 타당성을 찾는 데" 있으며, 문언의 통상적 의미에 충실한 해석을 원칙으로 하되 입법 취지와 목적, 제·개정 연혁, 법질서 전체와의 조화를 추가로 동원한다는 것이다.^[2]

이 한 문장 안에 네 개의 렌즈가 들어 있다.

1. **문리적 해석** — 용어와 문장의 통상적 의미로 읽음. 언제나 출발점임.

2. **체계적 해석** — 그 조문이 놓인 구조로 읽음. 같은 법의 다른 조문, 이웃한 법령과의 관계를 봄.
3. **역사적 해석** — 입법자가 그 조문을 만들 때의 의사로 읽음. 제·개정 이유서가 단서가 됨.
4. **목적론적 해석** — 법이 실현하려는 목표로 읽음. 신기술처럼 법문이 예상하지 못한 상황에서 무게가 커짐.

순서가 중요하다. 문언이 명확하면 거기서 멈춘다. 특히 국민의 권리를 제한하거나 의무를 지우는 조문일수록 엄격하게 읽는다. 렌즈를 겹칠수록 해석자의 재량이 커지고, 재량이 커질수록 같은 조문은 언제 누가 읽어도 같아야 한다는 요청이 흔들리기 때문이다.

같은 조문, 정반대의 결론

이론은 여기까지다. 실전은 더 흥미롭다.

산업통상자원부 법무팀의 내부 강의 자료에 실린 사례 하나.^[3] 자율주행차 실증특례(규제를 한시적으로 면제하고 실증하게 하는 제도)를 받은 기업이 특례 연장을 신청했다. 관련 법률은 개정이 끝났지만, 그 법을 실제로 작동시키는 고시는 아직 나오지 않은 상태였다. 연장 요건은 "법령이 정비되지 않은 경우"였다.

법률은 고쳐졌는데 고시가 없다면, 법령은 정비된 것인가 아닌가.

부처 자체 검토는 "고시·훈령까지 법령으로 본다. 고시가 없으니 미정비, 연장 가능"이라 읽었다. 외부 자문은 "법률 개정이 완료됐으니 정비 완료, 연장 불가"라 읽었다.

같은 조문이다. 같은 사실관계다. 결론은 정반대였다.

어느 쪽도 틀리지 않았다. 문언의 그물코가 넓은 자리에서 한쪽은 제도의 목적을, 다른 쪽은 문언의 경계를 무겁게 읽었을 뿐이다.

저울의 반대편도 있다. 오존층 보호 법령 위반을 인지한 공무원은 "고발하여야 한다." 문언은 의무다. 그런데 법원은 가별성이 없는 사정 등이 있으면 고발하지 않을 재량을 인정했다.^[4] "하여야 한다"라는 가장 단단해 보이는 문언조차 구체적 타당성 앞에서는 휘어진다.

[필자 분석] 이 오래된 문제가 지금 반복되고 있다

이 행정법 사례가 AI 규제의 한복판에서 그대로 재연되는 중이다.

한국 인공지능기본법은 프론티어 AI의 연산량 산정방식, 투명성 의무의 적용 예외, 공공 우선구매 대상의 확인 기준 같은 핵심 기준을 대부분 고시로 위임해 두었다. 고시가 나오기 전까지 법의 문장은 미완성이다. "법령이 정비되지 않은 경우"를 둘러싼 바로 그 해석 다툼이, 이번에는 AI 사업자의 의무 범위를 두고 벌어질 것이다.

EU도 다르지 않다. AI Act의 요건들은 조화표준과 가이드라인이 붙기 전까지 추상어에 가깝다.

규제의 실체는 법률의 층이 아니라 그 아래층에서 완성된다.

그래서 이 책은 하나의 규칙을 지킨다. 두 법을 비교할 때 "법령이 확정된 것"과 "고시를 기다리는 것"을 문장 단위로 구분해 적는다. 어느 쪽으로도 확인되지 않은 대목에는 [검토 필요]를 붙인다. 그 표시가 붙은 문장은 이 책의 논지가 기대는 자리에 두지 않았다.

이것이 이 책의 독법이다.

두 개의 법 앞에서

다음 장부터 우리는 한국 인공지능기본법과 EU AI Act를 나란히 놓는다. 그때 네 개의 렌즈를 들고 간다.

문언을 비교하는 것은 시작일 뿐이다. 두 법이 놓인 체계를 읽고 — 한국은 진흥법의 몸에 규제 조항을 엮었고, EU는 제품안전법의 뼈대 위에 서 있다 — 제정이유서에 남은 입법자의 의도를 읽고, 각 법이 실현하려는 목적을 읽어야 비로소 비교가 된다.

같은 법을 읽고도 결론이 갈리는 이유는 법이 부실해서가 아니다. 법이 언어이기 때문이다.

한줄 코멘트. 유럽의 40년 설계는 법이 기술을 따라가지 않도록 만든 구조였고, 그 대가로 법의 문장은 아래층이 채워지기 전까지 미완성으로 남는다. 두 AI법을 읽을 때 확정된 것과 기다리는 것을 갈라 두지 않으면, 비교는 시작하기도 전에 어긋난다.

공원의 안내판은 오늘도 서 있다. 유모차를 밀고 온 사람은 잠시 멈춰 서서 안내판을 읽는다. 그 잠깐의 멈춤에서 법은 시작된다.

주(註)

1. H. L. A. Hart, *The Concept of Law* (1961). "공원에 탈것 금지" 사례는 법문언의 핵심부와 반음영부를 설명하는 법철학의 고전적 예시다.
2. 대법원 2010. 12. 23. 선고 2010다81254 판결.
3. 산업통상자원부 법무TF, 「사례로 보는 산업부 법률 이야기 — 제1강: 법령해석 등」.
4. 서울고등법원 1970. 9. 3. 선고 69노558 판결(강의 자료 재인용).

PART 1

1부 — 설계도 대조: 두 법의 뼈대

같은 기술을 향한 두 개의 그물. 그물코가 왜 이렇게 다른가.

1장. 두 법은 다른 질문에서 태어났다

이 장을 한 문장으로

한국법은 "이 기술을 어떻게 키울 것인가"를 묻고, EU법은 "이 기술이 시장에 나와도 안전한가"를 묻는다. 뒤에 나올 모든 차이가 이 한 줄에서 갈라진다.

법에는 이름이 있다. 그리고 이름은 종종 그 법이 스스로를 어떻게 생각하는지를 가장 먼저 말해 준다.

2026년 7월 21일 한국에서 전면 시행된 법의 정식 명칭은 이렇다.

인공지능 발전과 신뢰 기반 조성 등에 관한 기본법^[1]

같은 해 8월 2일부터 EU에서 적용이 확대되는 법의 정식 명칭은 이렇다.

인공지능에 관한 조화규칙을 정하는 규정 (Artificial Intelligence Act)

[2]

두 문장을 나란히 놓으면 눈에 띄는 것이 있다.

한국 법의 이름은 "발전"이 맨 앞에 온다. EU 법은 "AI Act"조차 본제가 아니라 괄호 속 별칭이다. 본제는 "조화규칙을 정하는 규정" — 시장의 규칙을 통일하는 법이라고 스스로를 소개한다.

이름 하나로 다 알 수는 없다. 그러나 이 두 이름은 우연이 아니었다.

법에도 출생증명서가 있다

법률에는 제정이유서가 따라붙는다. 이 법을 왜 만드는지 국회와 정부가 스스로 밝히는 출생증명서다.

한국 인공지능기본법의 제정이유서는 이렇게 시작한다.

"인공지능의 건전한 발전을 지원하고 인공지능사회의 신뢰 기반 조성
필요한 기본적인 사항을 규정함으로써 국민의 권익과 존엄성을
보호하고 국민의 삶의 질 향상과 국가경쟁력을 강화하는 데 이바지..."^[3]

지원, 발전, 조성, 강화. 전부 무언가를 키우는 동사다.

EU AI Act 제1조는 이렇게 시작한다.

"본 규정의 목적은 역내시장의 기능을 개선하고, 인간 중심적이며
신뢰할 수 있는 인공지능의 도입을 촉진하는 한편, 건강·안전·기본권에
대한 높은 수준의 보호를 보장하는 것..."^[4]

촉진이라는 단어도 있다. 그러나 문장의 무게중심은 보장에 놓여 있다.

한 문장씩만 놓고 봐도 결이 다르다. 한국 법은 이 기술을 어떻게 잘 키울
것인가를 묻는다. EU 법은 이 기술이 시장에 나와도 안전한가를 묻는다.

조문을 세어 보면

말뿐이 아니다. 법의 몸통을 열어 보면 질문의 차이가 조문 배분으로 그대로
나타난다.

1. 한국 인공지능기본법 — 6장, 제43조까지. 제3장 산업 육성이 제13조에서 제26조까지, 제4장 윤리 및 신뢰성 확보가 제27조에서 제36조까지임. 조 번호로만 세면 14 대 10임.

2. 그런데 실제 조문 수는 46개임. 개정으로 붙은 가지번호 세 개 — 제17조의2(이용비용 지원), 제22조의2(인공지능연구소 설립), 제22조의3(연구기관 설립·운영) — 가 전부 산업 육성 장에 들어갔음. 실제 비율은 17 대 10으로 벌어짐. ^[5]
3. EU AI Act — 13장 113조. 혁신 지원(규제 샌드박스)을 다루는 장은 하나뿐이고, 나머지 열두 장은 전부 금지·고위험 요건·투명성·범용 모델 위험관리·거버넌스·시장감시·벌금임. 진흥은 부속 조항이고 본체가 규제임. ^[6]

세 가지를 겹치면 그림이 나온다. 한국법은 산업 육성이라는 몸통에 규제 조항을 얹은 구조다. EU법은 제품안전 규제라는 몸통에 AI라는 새 대상을 얹은 구조다.

2026년 1월 개정도 같은 방향이었다. 개정 이유서가 밝힌 무게중심은 전부 진흥과 지원과 포용이었고, 규제 조항인 제31조·제32조·제34조는 손대지 않았다.

무엇이 의무이고 무엇이 노력인가

여기서 흔한 오해 하나를 정리해야 한다. "한국법에는 강제 의무가 없다"는 말은 사실이 아니다.

한국법의 의무는 두 층으로 갈린다.

"이행하여야 한다" — 제31조 투명성 확보, 제32조 프론티어 모델 안전성 확보, 제34조 고영향 인공지능 사업자의 여섯 가지 책무, 제36조 국내대리인 지정. 이 넷은 의무다. 위반하면 사실조사와 시정명령이 따르고 과태료가 부과된다.

"노력하여야 한다" — 제30조제3항 고영향 AI의 사전 검·인증, 제35조 고영향 인공지능 영향평가. 이 둘은 노력 의무다.^[7]

의미심장한 것은 어느 쪽이 노력 의무로 남았는가이다. 사전에 제3자가 확인하는 절차 — 검·인증과 영향평가 — 가 정확히 그 자리에 있다.

EU는 반대다. AI Act 제43조는 고위험 AI에 적합성평가를 강제하고, 이를 거치지 않으면 시장에 낼 수 없다. 사전 관문 자체가 시장 진입의 조건이다.

한국법에도 의무는 있다. 다만 그 의무는 대체로 팔고 난 다음에 지켜야 할 것이고, EU의 의무는 팔기 전에 통과해야 할 것이다.

자백하는 제정이유서

흥미로운 지점이 있다. 이 구조가 은폐되어 있지 않다는 것이다. 오히려 제정이유서가 스스로 밝힌다.

한국법 제정이유서는 규제 조항조차 이렇게 설명한다.

"...민간이 자율적으로 추진하는 인공지능 안전성·신뢰성 검·인증 및 영향평가에 대한 지원을 하고..."^[8]

검·인증이 규제가 아니라 민간 자율 활동에 대한 지원으로 서술된다. 법이 스스로 규제 조항을 진흥의 언어로 번역해서 소개하는 셈이다.

여기서 단어 하나를 풀고 가자. 검·인증은 제30조제1항이 만든 법률상 정의어다.

인증은 이미 정해진 규격이 기준이다. "이 규격에 맞는가"를 묻는다.

검증은 그 규격이 없을 때 쓴다. 만든 쪽이 내건 주장을 기준 삼아 "당신 말이 맞는가"를 묻는다.

그래서 법이 굳이 둘을 붙여 썼다는 것 자체가 신호다. AI에는 아직 인증의 잣대가 될 규격이 없다.

한국에서 검증이 인증과 병렬로 놓였다는 것은, 기준이 아직 없다는 사실을 법이 문장 구조로 인정한 것이다. 유럽에서는 이 둘이 나란히 놓이지 않는다. 왜 그런지는 7장에서 조문으로 따라간다.

EU 쪽에도 비슷한 자기 확인이 있다. 2024년 5월 집행위원회 디지털총국 담당자는 웨비나에서 이렇게 말했다.

"AI Act는 NLF 체계의 일부다."^[9]

법 이름의 본체가 조화규칙이었던 이유가 여기서 확인된다. EU AI Act는 스스로를 새로운 종류의 법이 아니라 40년 된 제품안전 규제 체계의 최신 확장판으로 규정한다.

두 법 모두 겉으로 내세우는 이름과 속 구조가 정확히 일치한다. 숨기는 것이 아니라 자백하고 있다.

[필자 분석] 태생이 다른 두 법

앞으로 이 책이 다룰 모든 비교 — 고영향 대 고위험, 금지 목록의 유무, 검·인증의 강제성, 프론티어 모델의 문턱, 제재의 크기 — 는 사실 이 한 지점으로 수렴한다.

두 법은 서로 다른 질문에 대한 답이다.

한국법의 질문은 "이 기술을 어떻게 국가경쟁력으로 만들 것인가"다. 국가인공지능전략위원회가 대통령 소속인 것, 산업 육성이 법 제목 맨 앞에 오는 것, 검·인증이 노력 의무로 남은 것, 개정으로 늘어난 조문이 전부 진흥

장에 붙은 것 — 전부 이 질문의 논리적 귀결이다. AI는 이 법 안에서 일관되게 **키워야 할 자산**이다.

EU AI Act의 질문은 "이 기술이 시장에 나와도 시민이 안전한가"다. 0부에서 따라간 CE 마크와 적합성평가라는 40년 된 질문 형식에 AI라는 항목을 끼워 넣은 것이다. AI는 이 법 안에서 일관되게 **검증받아야 할 위험**이다.

어느 쪽이 옳은가를 묻는 것은 이 책의 목적이 아니다. 두 법은 각자의 질문 안에서 정합적이다.

문제는 이 출발점의 차이를 모른 채 두 법을 조문 대 조문으로 직역하려 할 때 생긴다. "고영향 인공지능"을 "고위험 AI"의 한국어 번역으로 읽는 순간, 산업 육성법의 개념을 제품안전법의 틀에 억지로 끼워 맞추는 오독이 시작된다.

그 오독이 실무에서 어떤 값을 치르게 하는지는 3장에서 본다.

질문이 다르면 같은 답도 다른 뜻이다

두 법 모두 "AI가 안전해야 한다"는 문장에는 동의할 것이다.

그러나 한국법에서 그 문장은 "국가경쟁력을 해치지 않는 선에서"라는 조건절을 품고 있고, EU법에서 그 문장은 "시장 진입의 전제조건으로서"라는 조건절을 품고 있다.

같은 문장, 다른 무게.

그렇다면 남은 질문은 하나다. 서로 다른 출생증명서를 줬 두 법이, 정확히 어느 지점에서 같은 대상을 다른 이름으로 부르고 있는가.

그 전에 확인할 것이 하나 있다. EU 쪽 출생증명서에 적힌 "NLF 체계의 일부"라는 자기 규정이 실제 조문에서 어떻게 확인되는지다.

실무자를 위한 체크박스

- ☐ 우리 서비스에 적용되는 한국법 조항이 "이행하여야 한다"인지 "노력하여야 한다"인지 구분했는가 (의무: 제31·32·34·36조 / 노력: 제30·35조)
- ☐ EU 쪽 의무 중 시장 진입 **전에** 통과해야 하는 것과 진입 **후**에 지켜야 하는 것을 갈라 두었는가
- ☐ 대외 메시지를 만들 때 한국법을 "규제법"이 아니라 "진흥법에 규제 조항이 없인 법"으로 설명하고 있는가 — 상대 부처의 언어와 어긋나면 대화가 걸돈다
- ☐ 사내에서 "고영향 = 고위험"으로 통용되고 있지 않은지 확인했는가

주(註)

1. 인공지능 발전과 신뢰 기반 조성 등에 관한 기본법(법률 제21311호, 2026.1.20 일부개정, 2026.7.21 시행). 원 제정은 법률 제20676호(2025.1.21).
2. Regulation (EU) 2024/1689, OJ L, 2024.7.12. "Artificial Intelligence Act"는 정식 명칭이 아니라 괄호 안 약칭으로 부기된다.
3. 인공지능기본법 원 제정(법률 제20676호) 제정이유, 법제처 국가법령정보센터 제공 원문.
4. EU AI Act Article 1 (Subject matter).
5. 법제처 국가법령정보 목차(2026-08-04 조회 기준) 확인. 제1장 총칙(제1~5조), 제2장 추진체계(제6~12조), 제3장 산업 육성(제13~26조 + 제17조의2·제22조의2·제22조의3), 제4장 윤리 및 신뢰성 확보(제27~36조), 제5장 보칙(제37~41조), 제6장 벌칙(제42~43조). 조 번호는 제43조까지이나 가지번호를 포함한 실제 조문 수는 46개다.
6. EU AI Act 13개 장 구성. 혁신 지원 조치는 Chapter VI(Article 57-63)에 한정된다.
7. 인공지능기본법 제30조제3항·제35조제1항("노력하여야 한다") 대 제31조·제32조·제34조·제36조("이행하여야 한다" 등 의무 규정). 위반 시 제40조 사실조사·시정명령과 제43조 과태료가 연결된다.
8. 인공지능기본법 원 제정(법률 제20676호) 제정이유 주요내용 (자)항.

9. 집행위원회 디지털총국·유럽 AI 사무국 웨비나(2024.5.30), "Risk management logic of the AI Act and related standards."

2장. AI Act는 새 법인가, 기존 법의 확장인가

이 장을 한 문장으로

AI Act 113개 조문 가운데 AI 고유의 논리로 새로 설계된 것은 넷뿐이고, 나머지는 40년 전 문법에서 가져왔다.

토스터와 챗봇은 닮은 데가 없다.

하나는 빵을 굽고 다른 하나는 말을 한다. 하나는 손에 잡히고 다른 하나는 화면 속에만 있다. 하나는 1920년대에 나왔고 다른 하나는 2020년대의 산물이다.

그런데 유럽에서 이 둘은 같은 체계 아래 놓인다. 토스터에 붙는 CE 마크가 고위험 AI 시스템에도 붙는다. 같은 적합성평가를 거치고, 같은 종류의 기술문서를 준비하고, 같은 시장감시를 받는다.

100년 차이가 나는 두 물건이 어떻게 하나의 규제 안에 들어가는가.

답은 하나다. EU는 AI를 위한 새 건물을 짓지 않았다. 40년 된 건물에 새 층을 올렸다.

정의 없이 등장하는 개념들

이 법을 꼼꼼히 읽으면 이상한 것을 발견한다. 핵심 개념들이 정의 없이 나온다.

"적합성평가 절차"가 Article 43에 나오는데 적합성평가가 무엇인지 이 법은 설명하지 않는다. "Notified Body"가 Article 33에 나오는데 그것이 무엇인지도 설명하지 않는다. "시장감시"가 Article 74에 나오는데 역시 정의가 없다.^[1]

빠뜨린 것이 아니다. 이미 다른 곳에 쓰여 있기 때문이다.

0부에서 따라간 그 체계다. 기초공사를 다시 할 필요가 없었다. 기초는 40년 동안 검증되었으니까.

다섯 개의 기둥은 그대로 서 있다

0부에서 확인한 설계 원리가 AI Act에 그대로 살아 있는지 대조해 보자.

기술 중립. "파라미터 몇 억 개 이상"이라고 쓰지 않는다. 대신 "정확성, 견고성, 사이버보안을 보장하라"고 쓴다.^[2] 구체적 수치는 법에 없다.

표준 위임. Article 40이 조화표준 준수 시의 준수 추정을 명문화한다. CEN-CENELEC JTC 21이라는 전담 위원회가 AI 조화표준을 개발하고 있다.^[3]

비례성. 모듈 A부터 H까지가 위험에 비례했듯, AI Act는 네 단계 피라미드로 이를 구현한다. 금지, 고위험, 제한적 리스크, 최소 리스크.

사전과 사후의 이중 구조. Article 43이 사전 적합성평가를, Article 74에서 79까지가 사후 시장감시를 규정한다.

경제적 운영자의 책임. 공급자, 수입업자, 배포자에게 각각의 의무가 부여된다.

다섯 기둥이 모두 서 있다. 건물이 같다.

새 층에서 달라진 네 가지

그러나 AI는 토스터가 아니다. 40년 된 건물에 새 층을 올리면서, 기존 설계로는 풀리지 않는 문제 넷이 드러났다.

배치자라는 새 행위자

기존 체계에서 제품이 시장에 나오면 그 뒤로 사용자는 규제 의무를 지지 않았다. 토스터를 산 사람에게 안전 의무가 없듯이.

AI는 다르다. 같은 시스템이라도 누가 어디에 어떻게 배치하느냐에 따라 위험이 완전히 달라진다. 채용에 쓰면 고위험이고 음악 추천에 쓰면 최소 리스크다.

그래서 AI Act는 배치자(Deployer)라는 새 규제 대상을 만들었다. 자신의 권한 아래 전문적 맥락에서 AI를 사용하는 자연인 또는 법인이다.^[4] 단순 소비자가 아니라 업무에 AI를 투입하는 조직이다. 40년 역사에 없던 개념이며, 배치자는 AI를 적절히 사용하고 인간 감독을 유지하고 사고를 보고해야 한다.

사후 모니터링의 의무화

기존 체계에서 사후 관리는 주로 시장감시기관의 몫이었다. 제조사가 자발적으로 하는 것은 있었지만 체계적 의무는 약했다.

AI Act Article 72는 다르다. 공급자가 직접 사후 모니터링 계획을 세우고 실행해야 한다. 의료기기 규정의 모델을 AI에 이식한 것이다. 중대 사고가 나면 15일 안에 보고해야 한다는 의무도 붙었다.

이유는 단순하다. 토스터는 출시 후 물리적으로 변하지 않지만, AI는 데이터가 바뀌고 사용 패턴이 바뀌고 성능이 변동한다.

기술문서의 확장

기존 기술문서에는 설계 사양, 제조 공정, 시험 결과가 담겼다.

AI Act의 기술문서는 훨씬 방대하다. 시스템의 일반 설명뿐 아니라 **훈련·검증·테스트 데이터의 세부 사항**, 사이버보안 조치, 로그 기록 역량, 리스크 관리 시스템 전체가 문서화되어야 한다.^[5]

토스터의 기술문서가 설계도와 시험 성적서라면, AI의 기술문서는 거기에 **훈련 과정의 일기장**까지 더한 셈이다.

지속적 준수 의무

기존 제품은 출시 시점에 요건을 충족하면 그 상태가 유지된다고 가정할 수 있었다. 회로는 스스로 바뀌지 않으니까.

AI는 학습하고 업데이트되고 환경에 적응한다. 그래서 AI Act는 **실질적 변경(Substantial Modification)**이라는 개념을 도입했다.^[6] 출시 후 AI가 크게 바뀌면 적합성평가를 다시 받아야 한다. 이 문제는 5장에서 따로 다룬다.

대응표

0부에서 다른 개념들이 AI Act의 어디에 대응하는지 한눈에 보자.

40년 된 문법	AI Act
필수요구사항	↔ Art.8-15 고위험 AI 요건
조화표준 + 준수 추정	↔ Art.40
적합성평가 모듈	↔ Art.43 적합성평가 절차
EU 적합성 선언	↔ Art.47
기술문서	↔ Art.11 + Annex IV
CE 마킹	↔ Art.49
수권대리인 / 수입업자 / 유통업자	↔ Art.22 / 23 / 24
검증기관(NB)	↔ Art.33-36
인정(Accreditation)	↔ Art.33 (Reg 765/2008 준용)
시장감시	↔ Art.74-79 (Reg 2019/1020 준용)
Safety Gate 통보	↔ Art.79 중대 리스크 통보
[기존 문법에 없음]	→ Art.26 배치자 - 신설
[기존 문법에 없음]	→ Art.72 사후모니터링 - 강화
[기존 문법에 없음]	→ Art.51-55 범용 AI 모델 - 신설
[기존 문법에 없음]	→ Art.64 AI Office - 신설

왼쪽 열은 0부의 요약이고, 오른쪽 열은 이 책 나머지의 지도다.

[필자 분석] 기초 없이 새 층만 올린다면

AI Act를 "세계 최초의 AI 법"이라 소개하는 것은 맞지만 절반의 진실이다.

113개 조문 가운데 순수하게 AI 고유 논리로 설계된 것은 생각보다 적다. 배치자, 범용 AI 모델, AI Office, 사후 모니터링 정도다. 나머지 — 적합성평가, CE 마킹, 기술문서, 검증기관, 시장감시, 경제적 운영자 의무 — 는 전부 기존 체계에서 가져온 것이다.

이 구별이 왜 중요한가.

EU가 AI Act를 빠르게 집행할 수 있는 이유는 법을 잘 써서가 아니라 40년간 기초를 다져 왔기 때문이다. 적합성평가를 수행할 기관이 이미

있고, 그 기관을 검증할 인정 체계가 이미 있고, 시장에서 문제를 잡아낼 감시망이 이미 있다. AI Act는 그 위에 대상 하나를 추가했을 뿐이다.

한국을 같은 잣대로 보면 그림이 달라진다.

한국 인공지능기본법에는 조화표준에 해당하는 장치가 없다. 적합성평가 모듈도, AI를 심사할 검증기관 체계도 없다. 시장감시를 총괄하는 일반법도 AI에 연결돼 있지 않다. 기초를 짓지 않은 상태에서 새 층부터 올린 형국이다.

그래서 한국법의 사전 관문이 노력 의무로 남은 것은 입법자의 의지 부족이라기보다 강제할 잣대가 아직 없기 때문이라고 읽는 편이 정확하다. 1장에서 본 "검·인증"이라는 합성어가 그 사실을 문장 구조로 인정하고 있었다.

여기서 실무적 함의가 나온다. 한국이 EU 수준의 AI 규제를 원한다면 손봐야 할 것은 AI법 조문이 아니라 그 아래층이다. 반대로 기업 입장에서는, EU 대응 과정에서 만든 기술문서와 품질 시스템이 한국의 잣대가 갖춰지는 날 그대로 자산이 된다.

남은 문

건물의 구조를 확인했다. 다섯 개의 기둥이 서 있고 네 개의 방이 새로 붙었다.

그런데 이 건물에서 가장 중요한 방이 아직 열리지 않았다. 무엇을 기준으로 "이 AI는 위험하다"고 판단하는가. 그리고 한국은 같은 판단을 어떤 이름으로 하고 있는가.

다음 장에서 그 문을 연다.

실무자를 위한 체크박스

- ☐ 우리가 EU 기준으로 공급자인지 배포자인지 판정했는가 — 배포자여도 인간 감독과 사고 보고 의무가 붙는다
- ☐ 기술문서에 훈련·검증·테스트 데이터의 세부 사항이 들어갈 준비가 되어 있는가 (Art.11 + Annex IV)
- ☐ 사후 모니터링 계획을 문서로 갖고 있는가. 중대 사고 15일 보고 절차를 정해 두었는가
- ☐ EU 대응으로 만든 산출물 중 한국 제34조 여섯 책무에 그대로 재사용할 수 있는 것을 정리했는가

주(註)

1. AI Act는 기존 체계의 개념을 정의 없이 참조한다. Blue Guide 2022(OJ C 247)가 이 개념들의 실질적 해설서 역할을 한다.
2. AI Act Article 15 (정확성·견고성·사이버보안).
3. AI Act Article 40. 최초로 공개심의에 오른 AI 조화표준은 prEN 18286(AI 품질경영시스템, 2025.10)이다. 조화표준 제정 진행 상황은 계속 바뀌므로 집필 시점에 재확인이 필요하다. [검토 필요]
4. AI Act Article 3(4) 및 Article 26.
5. AI Act Article 11 및 Annex IV.
6. AI Act Article 3(23). "시장 출시 또는 서비스 투입 후 AI 시스템의 변경으로서, 초기 적합성평가에 반영되지 않았거나 초기 적합성에 영향을 미칠 수 있는 것."

3장. 무엇을 위험이라 부르는가 — 고영향 vs 고위험

이 장을 한 문장으로

한국의 "고영향"과 EU의 "고위험"은 같은 자리에 있는 두 이름이 아니라, 서로 다른 나무에서 각자 자란 단어다.

응급실 문이 열리면 의사가 하는 첫 번째 일은 치료가 아니다.

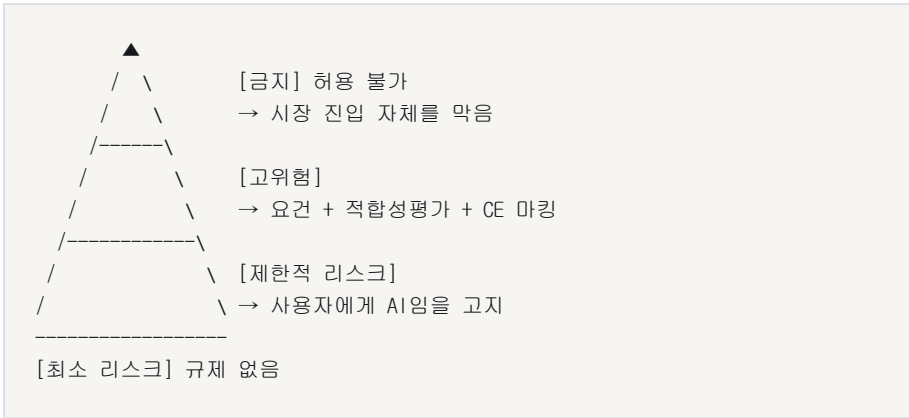
분류다.

이 환자는 지금 수술실로 가야 하는가. 저 환자는 30분 기다려도 되는가. 그 환자는 집에 가서 쉬면 낫는가. 의학에서 이를 트리아지라 부른다. 자원은 한정돼 있고 환자는 계속 온다. 모두를 같은 강도로 치료하는 것은 불가능하다.

EU AI Act는 AI에 대해 정확히 같은 일을 한다. 음악 추천 AI와 범죄 예측 AI를 같은 잣대로 다룰 수 없으니, 위험의 크기에 따라 분류하고 분류에 비례하는 규제를 매긴다. 0부에서 확인한 비례성이 여기서 작동한다.

EU의 피라미드

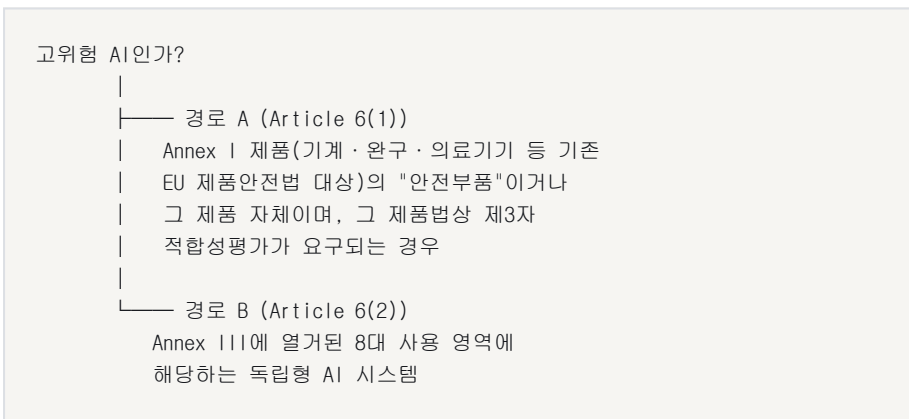
AI Act의 분류는 네 층 피라미드다.^[1]



바닥이 가장 넓다. 스팸 필터, 게임 AI, 재고 관리 AI 등 대부분의 AI가 여기 있고 아무 규제를 받지 않는다. 규제의 부피는 위로 갈수록 급격히 커진다.

고위험으로 가는 두 갈래

이 피라미드에서 가장 복잡한 층이 고위험이다. 그리고 여기에 도달하는 길이 둘이다.^[2]



경로 A는 **몸을 가진 AI**다. 이미 다른 법이 적용되고 있고 AI Act는 그 위에 올라탄다. 의료기기에 내장된 진단 AI, 자동차의 자율주행 AI, 기계의 안전 제어 AI가 여기다.

경로 B는 **몸이 없는 AI**다. 제품의 부품이 아니라 독립된 소프트웨어 자체가 판정 대상이다. Annex III의 여덟 영역은 이렇다.

영역	예시
생체인식	원격 생체인식 식별(금지 영역 제외)
핵심 인프라	도로교통 안전, 수도·가스·난방·전기
교육·직업훈련	입학 심사, 시험 평가, 부정행위 탐지
고용·인사관리	채용, 해고, 근무 배정, 성과 평가
필수 서비스 접근	대출 심사, 보험 가입, 긴급 서비스 배정
법 집행	증거 신뢰성 평가, 재범 예측, 프로파일링
이민·망명·국경	비자 심사, 허위 문서 탐지, 위험 평가
사법·민주주의	판결 지원, 법률 해석, 분쟁 해결

이 두 갈래는 낯선 구조가 아니다. 0부에서 본 조화 제품과 비조화 제품의 이원 구조와 같은 문법이다. EU는 AI 하나를 판정할 때도 이 오래된 두 갈래 습관을 반복한다.

고위험이 되면 무엇을 해야 하는가

고위험으로 확정되면 여덟 가지 기술 요건이 따라붙는다.^[3]

Art.9 리스크 관리 시스템 - 생애주기 전반의 지속 프로세스
 Art.10 데이터 거버넌스 - 훈련 데이터의 품질 · 대표성 · 편향
 Art.11 기술 문서화 - Annex IV, 설계부터 훈련 데이터까지
 Art.12 자동 로그 - 결정 과정을 추적 가능하게
 Art.13 투명성 · 정보 제공 - 배치자에게 사용 설명서
 Art.14 인간 감독 - 사람이 중단 · 무시할 수 있는 능력
 Art.15 정확성 · 견고성 · 사이버보안
 Art.17 품질 관리 시스템 - 위 전부를 관리하는 조직 체계

Article 9의 리스크 관리 시스템이 빠대다. 식별·분석 → 추정·평가 → 사후 데이터 기반 재평가 → 조치 채택의 네 단계를 계속 순환한다. 0부에서 본 3단계 저감 우선순위 — 설계로 제거, 보호 조치, 정보 제공 — 도 Article 9에 그대로 이식돼 있다.

한국의 한 줄 목록

한국 인공지능기본법 제2조제4호는 "고영향 인공지능"을 이렇게 정의한다.

사람의 생명·신체의 안전 및 기본권에 중대한 영향을 미치거나 위험을 초래할 우려가 있는 인공지능시스템

그리고 가목부터 차목까지 열 개 영역을 한 줄로 나열한다.^[4]

에너지 공급 / 먹는물 생산 공정 / 보건의료 제공·이용체계 / 의료기기 및 디지털의료기기 / 원자력시설 관리·운영 / 범죄수사·체포용 생체인식 분석 / 채용·대출심사 등 개인 권리·의무에 중대 영향 / 교통수단·시설·체계 / 국가기관등 의사결정 / 유아·초중등교육 학생 평가.

열 항목이 같은 층에, 같은 형식으로 나란히 앉아 있다. 나열에 위계가 없다.

겹쳐 보면 어긋난다

한국의 목록과 EU의 나무를 포개면 어긋나는 지점이 바로 드러난다.

한국 목록의 **의료기기**를 보자. 이것은 EU Annex III 여덟 영역 어디에도 없다. EU에서 의료기기에 내장된 AI는 경로 A로, 즉 의료기기 규정의 안전부품으로 분류된다. 같은 "의료 AI"를 놓고 한국은 하나의 목록 칸에 담았고 EU는 아예 다른 판정 나무의 다른 가지에 심어 놓았다.

교통수단도 마찬가지다. 자동차에 실린 AI는 EU에서 Annex III가 아니라 Annex I으로 가고, 그것도 CE 마킹 체계가 아닌 자동차 형식승인이라는 별도 경량체제로 흘러간다.^[5]

한국의 열 항목을 EU의 두 갈래 나무에 다시 심으면 대략 이렇게 갈린다.

1. **EU Annex III(경로 B)에 대응** — 채용·대출심사 → 고용·필수서비스, 범죄수사 생체인식 → 생체인식·법집행, 국가기관 의사결정 → 공공서비스·사법, 교육 → 교육.
2. **EU Annex I(경로 A)에 대응** — 의료기기 → 의료기기규정 안전부품, 교통 → 자동차·기계류, 원자력·에너지 → 인프라 관련 제품법.
3. **어느 쪽에도 직접 대응이 약한 것** — 먹는물 생산 공정, 그리고 개별 기기가 아니라 "체계"를 가리키는 보건의료 제공·이용체계.

하나의 한국어 단어가 EU의 서로 다른 두 세계에 걸쳐 있다.

확인하는 방식도 다르다

분류만 다른 것이 아니다. "내가 여기 해당하는지"를 확인하는 절차 자체가 다르다.

한국법 제33조는 사업자가 스스로 검토하고, 판단이 애매하면 과학기술정보통신부장관에게 확인을 요청할 수 있다고 정한다. 요청은 재량이다. 신청하지 않아도 법 위반이 아니다.^[6]

EU는 다르다. Annex III에 해당하면 고위험이 원칙적으로 **자동** 확정된다. 예외를 인정받으려면 네 유형 중 하나 — 좁은 절차작업, 이미 끝난 인적 활동의 개선, 패턴 편차 탐지, 예비 준비작업 — 에 해당해야 하고, 그 판단 근거를 문서로 만들어 보관해야 한다. 그리고 **프로파일링이 조금이라도 포함되면 이 예외 자체가 무효**가 된다.^[7]

고위험으로 확정되면 시장에 내놓기 전에 EU 데이터베이스에 공개 등록해야 한다. 등록하지 않은 고위험 AI는 애초에 시장에 나갈 수 없다.

한국에는 물어볼 수 있는 창구가 있고, EU에는 등록하지 않으면 열리지 않는 관문이 있다.

[필자 분석] 단어가 다른 이유

"영향"과 "위험"이라는 단어 선택은 우연이 아니다.

위험은 0부에서 추적한 개념이다. 확률과 심각도의 조합이며, 계산 가능한 크기다. 계산할 수 있기에 EU는 강제 등록과 제3자 평가와 CE 마킹이라는 사전 관문을 세울 수 있었다. 짚 수 있어야 문턱을 정한다.

영향은 계산의 언어가 아니라 **결과의 언어**다. "기본권에 중대한 영향을 미칠 우려"라는 문구는 확률을 계산하라고 요구하지 않는다. 결과가 클 수 있다는 사실만 확인하면 된다.

이 단어 선택은 1장에서 본 두 법의 출생증명서와 정확히 맞물린다. EU는 "시장에 내놓아도 안전한가"를 묻는 법이었기에 계산 가능한 위험 개념과 강제 관문이 필요했다. 한국은 "이 기술을 어떻게 키울 것인가"를 묻는

법이었기에, 사업자의 자기 판단에 기대는 영향 개념과 재량적 확인 절차로 충분했다.

목록의 항목 수가 열이나 여덟이나는 껍데기의 차이다. 진짜 차이는 목록 뒤에 무엇이 기다리고 있는가다.

요건의 밀도에서도 격차가 확인된다. EU의 여덟 요건 가운데 한국 제34조의 여섯 책무가 대응하는 것은 절반 남짓이다. 데이터 거버넌스, 자동 로그, 정확성·견고성·사이버보안에 해당하는 조항이 한국에는 없다.

여기서 실무적 경고가 양방향으로 나온다.

"우리 AI가 한국에서 고영향이 아니니 EU에서도 안전하겠지"라는 판단은 서로 다른 나무의 다른 가지를 같은 가지로 착각하는 것이다. 그리고 반대의 착각도 똑같이 위험하다. "EU Annex III에 없으니 한국에서도 고영향이 아니겠지"라는 판단은, 그 시스템이 실은 EU의 다른 나무에 걸려 있을 가능성을 놓친다.

번역이 아니라 대조가 필요한 이유가 여기 있다.

목록 앞에서 멈추지 않는 것들

열 개, 여덟 개. 숫자로만 보면 두 목록은 비슷한 크기다.

그러나 한국의 목록에는 아예 없는 것이 EU에는 있다. **고위험이 되기도 전에, 목록에 오를 자격조차 주지 않는 항목들이다.** 사회적 평점, 특정 취약계층을 겨냥한 조작, 무차별 얼굴 인식 데이터베이스. 이것들은 분류해서 관리할 대상이 아니라 애초에 시장에 존재해서는 안 되는 대상으로 취급된다.

목록에 오르는 것보다 더 근본적인 판정이 있다면, 그 선은 누가 어디에
긋는가.

실무자를 위한 체크박스

- ☐ 우리 시스템을 EU 기준으로 판정할 때 경로 A(제품 안전부품)와 경로 B(Annex III) 중 어느 쪽인지 확정했는가 — 둘은 적용 법령 자체가 다르다
- ☐ Article 6(3) 예외를 주장할 계획이라면, 그 근거를 문서로 남길 준비가 되어 있는가. 프로파일링이 조금이라도 들어가면 예외가 무효가 된다
- ☐ 고위험이면 EU 데이터베이스 등록이 시장 진입의 선행 조건임을 일정에 반영했는가
- ☐ 한국 제33조 확인 요청을 할 것인지 결정했는가 — 재량이지만, 하지 않으면 판정 책임이 전부 사내에 남는다
- ☐ 한국에 대응 조항이 없는 세 가지(데이터 거버넌스, 자동 로그, 정확성·견고성·사이버보안)를 EU 대응 과정에서 이미 갖추게 된다는 점을 사업 계획에 반영했는가

주(註)

1. AI Act Recital 26 및 집행위원회 디지털총국 웨비나(2024.5.30)의 리스크 기반 접근 설명.
2. AI Act Article 6(1)(경로 A) 및 Article 6(2)(경로 B).
3. AI Act Article 8-15 및 Article 17.
4. 인공지능 발전과 신뢰 기반 조성 등에 관한 기본법 제2조제4호 가목~차목. 카목은 대통령령 위임이며 현재 시행령에 추가 지정은 없다.
5. AI Act Annex I Section B(자동차 형식승인 관련 비-NLF 법령군) 및 Article 2(2).
6. 인공지능기본법 제33조(고영향 인공지능의 확인). 사업자의 사전 검토는 의무이나, 장관에 대한 확인 요청은 재량이다.
7. AI Act Article 6(3). 예외 네 유형과 프로파일링 수행 시 예외 무효. 등록 의무는 Article 49 및 Article 71.

4장. 아예 금지된 것 — 10개와 0개

이 장을 한 문장으로

금지 목록의 격차는 논의가 부족해서 생긴 것이 아니라, 진흥법이라는 법 형식이 처음부터 지니고 있던 한계다.

규제에는 여러 강도가 있다. 신고하게 하고, 등록하게 하고, 검사받게 하고, 허가받게 한다. 그런데 그 위에 한 칸이 더 있다.

만들지 말라고 하는 것.

이것은 다른 종류의 문장이다. 앞의 것들은 조건을 걸지만, 이것은 조건을 걸지 않는다. 기술이 불가능해서가 아니라 그 기술이 넘지 말아야 할 선을 넘기 때문에 막는다. 기술의 한계가 아니라 가치의 한계를 선언하는 문장이다.

EU는 그 문장을 썼다. 한국은 쓰지 않았다.

원래의 여덟

AI Act Article 5는 여덟 가지를 금지했다.^[1] 2025년 2월 2일부터 이미 시행 중이다.

1. **잠재의식 조작** — 사람이 인식하지 못하는 방식으로 행동을 조작하는 AI임.

2. **취약성 악용** — 나이·장애·사회경제적 상황 같은 취약점을 이용해 행동을 왜곡하는 AI임.
3. **사회적 평점** — 사회적 행동이나 성격 특성에 점수를 매겨 불이익을 주는 AI임.
4. **개인 범죄 리스크 예측** — 프로파일링만으로 특정 개인의 범죄 가능성을 예측하는 AI임.
5. **무차별 얼굴 데이터베이스 구축** — 인터넷이나 CCTV에서 무차별로 얼굴 이미지를 긁어모아 DB를 만드는 AI임.
6. **직장·학교 감정인식** — 업무 현장이나 교육기관에서 감정을 추론하는 AI임. 의료·안전 목적은 예외임.
7. **민감 특성 기반 생체 분류** — 인종·정치 성향·종교 같은 민감 속성으로 사람을 분류하는 AI임.
8. **공공장소 실시간 원격 생체인식** — 법 집행 목적의 실시간 얼굴인식임. 실종아동 수색이나 임박한 테러 위협 등 극히 제한된 예외만 허용되며 사법부의 사전 승인이 필요함.

여덟 항목을 관통하는 것은 하나다. 이 AI들이 겨냥하는 것은 사람의 지갑이나 신체가 아니라 **인간이 스스로를 결정할 수 있다는 전제 자체**다.

새로 온 들

숫자가 바뀌었다. 지금은 여덟이 아니라 열이다.

2026년 협상 과정에서 두 항목이 새로 끼어들었다.^[2]

- 아동 성적 학대 소재(CSAM)를 생성하는 AI

- 비동의 친밀 이미지(NCII)를 생성하는 AI — 이른바 "누디파이어" 앱을 포함함

법 조문 하나가 협상 테이블 위에서 몇 달 만에 두 줄 늘어난 것이다.

여기서 눈여겨볼 것이 있다. 이 두 금지는 "그런 이미지를 만들 수 있는 AI는 전부 불법"이라는 무차별 금지가 **아니다**. 조건부 이원 테스트를 통과해야 금지에 걸린다.^[3]

모델을 만든 쪽(공급자) 기준으로는 둘 중 하나여야 한다. 그런 생성이 시스템의 **의도된 목적**이거나, 설계·훈련상 그런 결과가 **합리적으로 예견** 가능한데도 적절한 안전조치를 두지 않은 경우다. 범용 이미지 생성 AI가 이론적으로 악용될 수 있다는 사실만으로는 걸리지 않는다. 관건은 **세이프가드의 유무**다.

실제로 쓰는 쪽(배포자) 기준은 더 단순하다. 실제로 그 목적으로 사용한 경우만 해당하고, 우연한 생성이나 다른 합법적 목적의 사용은 제외된다.

정교한 설계다. "이런 것을 만들 수 있는 기술 자체를 없애라"가 아니라 "안전장치 없이 방치하지 마라, 그리고 실제로 악용하지 마라"는 이중의 조건문이다.

이 두 항목은 이제 확정됐다. **Regulation (EU) 2026/1744**로 2026년 7월 8일 최종 채택되고 관보에 게재됐으며, 적용일은 2026년 12월 2일이다.

한국의 숫자는 0이다

한국 인공지능기본법에는 이런 조문 자체가 없다. 여덟이 열이 되는 동안 한국의 숫자는 계속 0이었고 지금도 0이다.

여기서 성급한 결론 하나를 건어내야 한다. "한국은 0개다"라는 문장이 "그러니 한국에서 이런 것들이 합법이다"를 뜻하지는 않는다. 인공지능기본법 **바깥**의 일반 형사법이 이런 콘텐츠의 제작과 유포를 별도로 처벌할 가능성이 높다.^[4]

그러나 **바깥에** 있다는 사실 자체가 중요하다.

EU는 이 금지를 AI 규제법 안에 넣었다. 그래서 AI Act 위반이 되고, AI Office와 시장감시당국이 감독하고, 매출 대비 과징금이라는 AI Act의 집행 도구가 그대로 작동한다. 모델 공급자에게 사전 안전조치 의무까지 얹혔다.

한국이 이 문제를 다루는 곳은 형법 계통의 일반 형사법이다. 사후 처벌은 가능하다. 그러나 안전조치를 갖추지 않은 AI 시스템 자체를 사전에 겨냥하는 도구는 아니다.

같은 목적을 다른 층에서 다룬다. 그리고 그 층의 차이가 실효성의 차이로 이어진다.

[필자 분석] 금지는 진흥법의 문법이 아니다

이 공백을 두고 흔히 "논의가 부족했다"거나 "입법이 늦었다"고 말한다. 앞의 세 장을 따라온 지금은 더 정확한 답을 낼 수 있다.

1장에서 확인했다. 한국 인공지능기본법은 "이 기술을 어떻게 키울 것인가"라는 질문에서 태어났다. 3장에서 확인했다. 그 질문은 계산 가능한 "위험"이 아니라 결과 중심의 "영향"이라는 언어를 낳았다.

금지라는 규제 수단은 이 언어 위에서 애초에 어색하다.

금지는 "이것은 시장에 존재해서는 안 된다"고 선언하는 문법이다. 진흥법의 기본 동사는 지원하다, 조성하다, 육성하다다. 그 동사들 사이에 "금지한다"를 끼워 넣는 순간 법 전체의 정체성이 흔들린다. 산업 육성을 다루는 열일곱 개 조문 옆에 "다음은 만들면 안 된다"는 조항이 들어설 자리는 원래 좁다.

반대로 EU AI Act는 제품안전법의 계보 위에 있다. 그 계보는 "위험한 제품은 시장에서 배제한다"는 문법을 40년 넘게 써 왔다. Article 5는 그 오래된 배제의 문법을 AI에 적용한 것뿐이다. EU에게 금지는 낯선 도구가 아니라 원래 서랍에 있던 도구다.

그러니 이것은 논의의 부재만은 아니다. 법의 몸통이 산업 육성법이면, 금지 조항은 이질적인 장기다.

한국이 이 공백을 메우려면 길은 둘이다. 인공지능기본법 안에 이질적인 규제 장기를 이식하거나, 형법 계통의 개별 입법으로 사안별 대응을 이어가거나. 지금까지의 선택은 후자였다. 그 선택의 대가는 분명하다. 이미 만들어진 것을 처벌할 수는 있어도, 만들어지기 전에 막을 도구는 갖지 못한다.

숫자 뒤의 질문

여덟에서 열로. 숫자는 계속 움직일 것이다.

그러나 정말 중요한 것은 숫자가 아니다. 어떤 법 형식이 금지라는 도구를 자연스럽고 품고, 어떤 법 형식이 그 도구를 이질적으로 느끼는가다. 숫자를 늘리기 전에 물어야 할 것은 법의 몸통이다.

만드는 것 자체를 막는 문제가 이렇다면, 만들어진 것이 계속 변해 갈 때는 어떻게 되는가.

실무자를 위한 체크박스

- ☐ 우리 모델이 이미지·영상 생성 기능을 포함한다면, CSAM·NCII 관련
세이프가드를 **설계 단계에서** 갖추고 그 사실을 문서로 남겼는가 — 이원
테스트에서 갈리는 지점이 정확히 여기서다
- ☐ EU에서 금지 목록에 걸리는 기능이 우리 제품에 부수적으로라도 들어 있지
않은지 점검했는가 (특히 감정인식, 생체 분류, 얼굴 데이터 수집)
- ☐ 한국 사업만 하더라도, 금지 목록에 해당하는 기능은 인공지능기본법이 아니라
형사법 리스크로 별도 검토했는가
- ☐ 금지 목록은 개정으로 늘어났다. 목록 변동을 모니터링 항목에 넣어 두었는가
(14장·부록 B 참조)

주(註)

1. AI Act Article 5. 여덟 가지 금지 관행은 2025년 2월 2일부터 적용된다.
2. Regulation (EU) 2026/1744(Digital Omnibus on AI) Article 1(7), AI Act Article 5(1) (ba)·(bb) 신설. 2026-07-08 채택, 관보 게재 2026-07-24, 발효 2026-07-27. 적용일은 개정 AI Act 제113조에 따라 2026-12-02.
3. Digital Omnibus, Article 5(1a) 신설. 공급자 기준(의도된 목적 또는 예견 가능성 + 안전조치 부재)과 배포자 기준(실제 목적 사용)의 이원 테스트. 밝기·배경 조정처럼 노출도를 늘리지 않는 단순 보정은 "조작"에 해당하지 않는다.
4. 한국 형사법 계통의 딥페이크·아동 성착취물 관련 처벌 조항과의 정확한 대응 관계는 원문 조문 대조를 거치지 않았다. 여기서는 "인공지능기본법 안에 금지 조항이 없다"는 사실과 "일반 형사법의 규율 가능성"을 구조적으로만 구분했다. 구체적 조문과 처벌 범위는 별도 확인이 필요하다. [검토 필요]

5장. 살아 움직이는 제품 — 업데이트와 재인증

이 장을 한 문장으로

40년 된 인증 체계는 "제품은 팔린 뒤 변하지 않는다"는 전제 위에 서 있었고, AI가 그 전제를 깼다.

작년에 산 스마트폰을 생각해 보자.

오늘 아침 운영체제가 업데이트됐다. 카메라 앱이 바뀌었고 보안 패치가 적용됐고 AI 비서의 성능이 달라졌다. 지난달에도 업데이트가 있었고 그 전달에도 있었다.

질문 하나. 오늘 당신이 쓰는 폰은 작년에 인증받은 그 폰인가.

물리적으로는 같은 하드웨어다. 그러나 소프트웨어는 수십 번 바뀌었다. 만약 이것이 의료기기였다면. 자율주행차의 두뇌였다면. 채용 지원자를 평가하는 시스템이었다면.

인증받은 제품이 인증받지 않은 제품으로 서서히 변해 가는 과정. 이것이 40년 체계가 마주한 가장 까다로운 문제다.

토스터는 변하지 않는다

0부에서 본 체계는 한 가지를 전제로 설계됐다. 제품은 출시 후 변하지 않는다.

적합성평가도, 검증기관 심사도, 리스크 평가도 전부 출시 시점의 스냅샷을 기준으로 작동한다. 공장에서 나오는 순간 설계대로 만들어졌는지 확인하고, 확인이 끝나면 CE 마크를 붙여 내보낸다.

토스터의 발열 코일은 스스로 재설계되지 않는다. 계량기의 측정 알고리즘은 밤사이에 업데이트되지 않는다. 물리적 제품은 시간이 지나도 같은 제품이다.

AI가 이 전제를 깨뜨린다.

AI는 세 가지 방식으로 변한다

의도적 업데이트. 개발자가 모델을 개선하거나 버그를 고친다. 새 데이터로 재훈련하고 무선으로 배포한다. 스마트폰 업데이트와 같은 방식이다.

지속적 학습. 일부 AI는 운영 중에도 새 데이터에서 계속 배운다. 사용자 행동에 적응하고 피드백으로 모델을 조정한다. 개발자가 아무것도 하지 않아도 AI가 스스로 변한다.

데이터 드리프트. 모델 자체는 그대로인데 입력 데이터의 성격이 바뀌면 출력이 달라진다. 팬데믹 기간에 대출 심사 AI가 오작동한 사례가 대표적이다. 훈련 데이터에 없던 패턴이 현실에 나타나면, 같은 모델이 다른 답을 낸다.

셋 다 같은 결과로 이어진다. 출시 시점에 적합성평가를 통과한 AI가 시간이 지나면 적합성을 잃을 수 있다.

실질적 변경이라는 해법

EU AI Act는 실질적 변경이라는 개념으로 답했다.^[1]

시장 출시 또는 서비스 투입 후 AI 시스템에 가해진 변경으로서, 초기 적합성평가에서 반영되지 않았거나 초기 적합성평가 결과에 영향을 미칠 수 있는 것

둘 중 하나만 해당해도 실질적 변경이다. 그리고 실질적 변경이 발생하면 적합성평가를 처음부터 다시 받아야 한다. CE 마킹도 다시, 기술문서도 갱신이다.

건물을 준공 검사받은 뒤 벽지를 바꾸는 것은 문제가 아니지만, 내력벽을 철거하면 다시 검사를 받아야 하는 것과 같다. 문제는 AI에서 내력벽과 벽지의 구분이 선명하지 않다는 것이다.

여기에 규정 하나가 더 붙는다. 이미 시장에 있는 고위험 AI에 실질적 변경을 가한 자는, **설령 원래 공급자가 아니더라도 새로운 공급자로 간주된다.**^[2] 2장에서 본 배치자와 공급자의 경계가 수정 행위 하나로 넘어갈 수 있다는 뜻이다. AI를 사다가 크게 손보면, 사용하는 쪽이 아니라 만든 쪽의 의무를 전부 떠안는다.

변하지 않아도 감시한다

실질적 변경이 없더라도 공급자에게는 지속적 감시 의무가 있다.

Article 72의 사후 모니터링이다. 고위험 AI를 시장에 낸 뒤에도 모니터링 계획을 세워 실행하고, 성능과 준수 상태를 능동적으로 수집·분석하고, 필요하면 시정조치를 해야 한다. 의료기기 규정의 모델을 그대로 가져온 것이다.

중대 사고를 인지하면 즉시, **늦어도 15일 안에** 해당 회원국 시장감시기관에 보고해야 한다.^[3] 0부에서 본 경보망이 AI에도 그대로 연결된다.

건강검진과 주치의

여기서 조문 제목 하나를 눈여겨볼 필요가 있다.

Article 9의 제목은 "리스크 관리 시스템"이다. "리스크 평가"가 아니다. 이 단어 선택은 우연이 아니다.

국제표준은 둘을 명확히 나눈다. 리스크 관리의 전체 과정은 맥락 설정 → 평가(식별·분석·판정) → 처리 → 모니터링과 기록의 순환이다. 즉 **평가는 전체 그림의 한 조각이고, 관리는 그 조각을 포함한 전체 사이클이다.**^[4]

건강검진을 생각하면 쉽다. 하루를 비우고 피를 뽑고 사진을 찍으면 며칠 뒤 결과지가 온다. 그 종이 한 장이 판정이다. 그런데 만성질환자에게는 검진 한 장으로 끝나지 않는다. 주치의가 매달 수치를 재고 약을 조정하고 이상 징후에 즉시 개입한다.

검진이 **한 번의 판정**이라면 주치의는 **계속되는 관계**다.

흥미로운 것은 EU가 모든 제품에 주치의 모델을 쓰지는 않는다는 점이다. 기계류 같은 물리적 제품에 요구되는 리스크 평가는 설계·출시 전 단계에 집중되고, 그 반복은 시장에 나오기 전에 종결된다. 일반 제품안전 쪽의 평가 방법론도 당국이 위험 제품을 걸러내는 분류 도구이지 기업에게 상시 체계를 요구하는 것이 아니다.

물리적 제품에는 검진을, AI에는 주치의를 붙였다. **팔린 뒤에도 토스터는 토스터이지만 AI는 아니기 때문이다.**

이 구분이 실무에서 갖는 무게는 작지 않다. 평가를 요구받은 기업은 정해진 시점에 정해진 형식으로 보고서 한 장을 준비하면 된다. 관리 시스템을 요구받은 기업은 조직 안에 상시 작동하는 절차를 만들고, 계속 기록하고, 이상 징후에 반응할 인력과 프로세스를 갖춰야 한다.

전자는 프로젝트고 후자는 조직 기능이다.

아직 답이 없는 것들

솔직하게 짚고 가야 할 것이 있다. EU의 해법도 완결되지 않았다.

고위험 채용 AI가 매주 새 이력서 데이터로 미세 조정된다면 매주 적합성평가를 다시 받아야 하는가. 어디까지가 실질적이고 어디까지가 미세한가.

지속학습 AI의 성능이 점진적으로 변할 때, 어느 시점에서 실질적 변경이 발생했다고 판단하는가. 1퍼센트일 때인가, 5퍼센트일 때인가.

범용 AI 모델이 업데이트되면 그 위에서 도는 수천 개의 하류 시스템이 모두 영향을 받는다. 전부 재평가해야 하는가.

AI Act는 아직 여기에 답을 내놓지 않았다. **실질적 변경의 구체적 기준은 조화표준과 가이드라인으로 채워질 영역이다.** 0-3에서 확인한 그대로다. 규제의 실체는 법률의 층이 아니라 그 아래층에서 완성된다. [검토 필요]

[필자 분석] 인증의 내구성

40년 체계의 핵심 강점은 "한 번 설계하면 오래 가는 법"이었다. 기술 중립성 덕분에 법을 자주 고칠 필요가 없었다.

그런데 AI에서는 법이 아니라 **제품 자체가 계속 변한다.** 시험대에 오른 것은 법의 내구성이 아니라 **인증의 내구성**이다.

한국은 이 문제에 더 취약하다. 인공지능기본법에는 **실질적 변경이라는 개념 자체가 없다.** 사후 모니터링 의무도 없고 중대 사고 보고 의무도 없다. AI가 출시 후 어떻게 변하든 법적으로 추적하거나 재평가할 장치가 마련돼

있지 않다.^[5] 있는 것은 제40조의 사실조사권뿐이며, 이는 문제가 이미 드러난 뒤에 작동하는 도구다.

그리고 이 공백은 3장에서 본 언어 선택과 정확히 이어진다. 한국법이 요구하는 것은 영향평가다. 평가는 시점의 언어다. 시점의 언어로는 계속 변하는 대상을 붙잡을 수 없다.

AI 규제의 진짜 어려움은 최초 인증이 아니라 지속적 준수에 있다. 이 지점에서 EU는 불완전하나마 체계를 갖추기 시작했고, 한국은 아직 출발선에 서지 못했다.

1부를 마치며

여기까지가 두 법의 뼈대다.

다른 질문에서 태어났고(1장), 한쪽은 40년 된 문법을 물려받았고(2장), 무엇을 위험이라 부를지에서 갈렸고(3장), 아예 만들지 말라고 할 수 있는지에서 다시 갈렸고(4장), 팔린 뒤의 변화를 붙잡는 방식에서 또 갈렸다(5장).

다섯 갈래가 전부 1장의 한 줄에서 뻗어 나왔다. 질문이 다르면 답도 다르다.

한줄 코멘트. "평가하라"와 "시스템을 관리하라"는 같은 문장처럼 보이지만 전혀 다른 약속이다 — 앞은 보고서 한 장이고 뒤는 조직 하나다. 두 AI법의 격차가 실무 비용으로 환산되는 첫 지점이 여기이며, EU 시장에 나가는 순간 그 격차는 준비 기간의 차이로 돌아온다.

뼈대를 봤으니 이제 살을 볼 차례다. 무엇을 표시하고, 무엇을 인증받고, 어디서부터 감시받고, 여기면 얼마를 무는가. 2부에서 돈과 시간이 실제로 드는 지점들을 하나씩 연다.

실무자를 위한 체크박스

- ☐ 모델 업데이트 이력을 **버전 단위로 기록**하고 있는가. 실질적 변경 판정은 결국 이 기록 위에서 이뤄진다
- ☐ 어느 정도의 변경을 실질적으로 볼지 사내 기준을 정해 두었는가. 기준이 없으면 판정 시점마다 다른 답이 나온다
- ☐ 남의 AI를 가져다 손보고 있다면, 그 수정으로 우리가 **공급자가 되는지** 확인했는가 (Art.25)
- ☐ 사후 모니터링이 "보고서 담당자 한 명"이 아니라 **상시 프로세스**로 설계돼 있는가
- ☐ 중대 사고 인지에서 보고까지 15일 안에 도는 사내 절차가 실제로 존재하는가 — 인지 시점을 누가 판단하는지까지 정해야 한다

주(註)

1. AI Act Article 3(23).
2. AI Act Article 25. 실질적 변경을 가한 자, 또는 상표를 바꿔 자기 명의로 출시한 자는 공급자로 간주된다.
3. AI Act Article 72(사후 모니터링) 및 Article 73(중대 사고 보고). 모델이 된 것은 의료기기 규정 Regulation (EU) 2017/745의 시판후조사 체계다.
4. ISO 31000:2018 Section 6의 순환 구조와 ISO 14971:2019의 절 구성(제4절 리스크 관리 시스템 / 제5-6절 분석·평가 / 제10절 사후생산 활동). AI Act Article 9는 시스템 구축·문서화·유지 의무(para.1), 4단계 프로세스(para.2), 통제 계층(para.5), 테스트 의무(para.6-8)로 구성된다.
5. 인공지능기본법에는 실질적 변경·사후 모니터링·중대 사고 보고에 해당하는 개념이 없다. 제40조 사실조사권이 사후 도구로 존재한다.

PART 2

2부 — 의무의 실제: 무엇을 해야 하는가

돈과 시간이 실제로 드는 지점들.

6장. AI가 만든 것임을 밝히는 두 방식

이 장을 한 문장으로

두 법 모두 "AI가 만든 것임을 밝히라"고 명령한다. 그런데 하나는 사람의 귀를 향하고, 다른 하나는 기계의 눈을 향한다.

밤 열한 시, 방의 불은 꺼져 있고 화면만 밝다.

대화창에 고민을 적는다. 답이 온다. 따뜻하고, 빠르고, 지치지 않는 답이.

그 답을 쓴 것은 사람이 아니다.

몸이 없으면 그물에 걸리지 않는다

0부에서 본 40년 설계는 한결같은 문법 위에 서 있었다. 몸이 무거울수록 그물은 촘촘해진다.

그런데 이 대화창 속 AI에는 몸이 없다. 완구가 아니고 전기제품이 아니고 기계가 아니다. 의료 목적을 표방하지 않는 한 의료기기도 아니다. 뒤집어서 CE 마크를 확인하고 싶어도 뒤집을 바닥 자체가 없다.

몸을 기준으로 짜인 그물은 몸이 없는 것을 거를 수 없다.

그럼 이 AI는 무법지대에 사는가. 아니다. EU의 답은 문법을 바꾸는 것이었다.

물건을 검사할 수 없다면, 말 그 자체에 표시를 붙인다.

다섯 항목, 세 개의 대응

AI Act Article 50이 그 조항이다. 다섯 개 항목으로 나뉜다.^[1] 한국 제31조와 하나씩 맞대 보자.

1. **Art.50(1)** — AI와 대화 중임을 이용자가 알 수 있게 설계할 것. "저는 AI입니다"가 설계 요건이 됨. 한국 제31조제1항의 사전 고지 의무가 여기 대응함.
2. **Art.50(2)** — AI 생성 콘텐츠에 기계판독 가능한 표시를 심을 것. 한국 제31조제2항의 생성물 표시 의무가 대응하나, 뒤에서 보듯 대응의 깊이가 다름.
3. **Art.50(3)** — 감정인식·생체분류 AI는 당사자에게 고지할 것. **한국에는 대응 조문이 없음.**
4. **Art.50(4)** — 딥페이크는 AI 생성·조작 사실을 명확히 고지할 것. 한국 제31조제3항이 대응함.
5. **Art.50(5)** — 공익 관련 AI 생성 텍스트는 AI가 썼다는 사실을 고지할 것. **한국에는 대응 조문이 없음.**

셋은 대응하고 둘은 아예 없다. 그런데 진짜 이야기는 "대응하는" 셋 안에 숨어 있다.

종이 태그와 RFID 칩

한국 제31조제2항은 생성형 AI 결과물임을 표시하라고 정한다. 시행령이 제시하는 표시 방법은 둘인데, **하나를 골라도 되고 안내 문구만으로도 된다.**^[2]

Art.50(2)는 다르다. 기계판독 가능한(machine-readable) 형식으로 표시하라고 못 박는다. 사람이 눈으로 읽는 안내 문구가 아니라 소프트웨어가 자동으로 인식할 수 있는 신호여야 한다. 워터마크, 메타데이터, 디지털서명 같은 기술이 이 요건을 채운다.^[3]

옷을 생각하면 쉽다.

"이 옷은 면 100%입니다"라는 종이 태그를 다는 것과, 옷감 안에 RFID 칩을 심어 스캐너를 대면 성분이 자동으로 뜨게 만드는 것. 둘 다 성분을 표시하는 행위다. 그러나 사람이 태그를 직접 읽어야 하는 것과 기계가 알아서 읽어내는 것은 완전히 다른 종류의 표시다.

콘텐츠는 그 자리에 머물지 않는다

왜 이 차이가 중요한가.

콘텐츠가 처음 만들어진 자리에만 머무른다면 안내 문구로 충분하다. 사람이 그 화면을 보고 "아, AI가 만든 거구나"를 안다.

그러나 콘텐츠는 머무르지 않는다. 캡처되고, 잘리고, 다른 플랫폼에 재업로드되고, 또 다른 AI의 학습 데이터로 흘러 들어간다.

그리고 이 이동의 대부분을 처리하는 것은 사람이 아니라 기계다.

플랫폼의 필터링 알고리즘, 검색엔진의 색인 시스템, 다음 세대 모델의 데이터 파이프라인. 이들은 화면에 적힌 안내 문구를 읽지 못한다. 파일 안에 박힌 신호만 읽는다.

안내 문구는 콘텐츠가 원래 자리에 있을 때만 작동한다. 기계판독 신호는 콘텐츠가 어디로 옮겨가도 따라붙는다.

둘의 실효 범위는 처음부터 다르게 설계됐다.

적용 시점에도 함정이 있다. Art.50(2)는 2026년 8월 2일 이후 새로 출시되는 시스템에 즉시 적용된다. 그 전에 이미 출시된 시스템만 2026년 12월 2일까지 경과조치를 받는다.^[4] "4개월 미뤄졌다"로 읽으면 오독이다. 신규 시스템은 유예 없이 정면으로 맞는다.

대화창 뒤의 심장

한 겹이 더 있다.

화면 속 앱은 껍데기이고, 답을 만드는 심장은 그 뒤의 **범용 AI 모델(GPAI)**이다. AI Act는 2025년 8월부터 이 모델을 만드는 회사를 직접 규율한다.^[5]

모델이 어떻게 학습됐는지 기술문서를 만들 것. 모델을 받아 쓰는 쪽에 정보를 제공할 것. 저작권 정책을 갖추고 저작권자의 기계판독 유보 표시를 존중할 것. 학습 데이터의 요약을 공개할 것.

물건의 세계에서 제조사가 지던 의무를, 말의 세계에서는 **모델 제공자가** 진다. 몸이 없어도 만든 손은 있기 때문이다.

몸을 요구하지 않는 것이 둘 더 있다

표시와 고지 말고도, 몸의 유무와 무관하게 작동하는 장치가 둘 있다.

금지선이 그렇다. 4장에서 본 목록은 물건이나 소프트웨어냐를 묻지 않는다. 사람의 취약한 마음을 이용해 행동을 조작하도록 설계된 AI라면, 인형에 들어 있는 앱 안에 있는 시장 진입 자체가 막힌다. 그리고 이때의 "중대한 피해"에는 신체만이 아니라 **심리적 건강**이 명문으로 들어 있다.^[6]

배상도 몸을 요구하지 않게 됐다. 개정된 제조물책임지침은 "제품"의 정의에 소프트웨어를 통째로 넣었다. 앱도 제품이고, 그 앱이 입힌

의학적으로 인정된 정신건강 손상도 배상 대상이며, AI의 복잡성 때문에 피해자가 결함을 입증하기 어려우면 법원이 추정으로 다리를 놓아 줄 수 있다.^[7] 이 이야기는 10장에서 본격적으로 다룬다.

여기에 아이러니가 하나 있다.

같은 대화 AI라도 **인형에 넣으면** 완구 관련 법과 무선기기 관련 법이 최소한 몸과 데이터는 지켜 준다. 그러나 그 인형이 하는 **말에** 대한 규율은 한참 뒤에야 도착한다.

그런데 **앱 속의** 같은 AI에게는 "너는 AI임을 밝혀라"는 의무가 곧바로 적용된다.

몸이 없는 쪽의 말이 몸이 있는 쪽의 말보다 먼저 규율되는 셈이다. 규제는 위험의 크기 순서가 아니라, 법이 붙잡을 수 있는 손잡이의 순서로 도착한다.

[필자 분석] 누구에게 말을 거는 법인가

1장에서 본 두 법의 출생증명서, 3장에서 본 "영향"과 "위험"의 언어 차이가 여기서 다시 나타난다.

한국 제31조는 사람에게 말을 거는 조문이다. AI가 결정에 관여했다는 사실을 알린다는 것은 사람과 사람 사이의 신뢰 행위다. 진흥법의 문법 안에서 투명성은 "정직하게 밝히기"로 완성된다.

Art.50(2)는 기계에게도 말을 거는 조문이다. EU는 콘텐츠가 인간의 눈을 거치지 않고도 자동화된 생태계를 관통한다는 사실을 전제한다. 그래서 사람이 읽는 안내문과 별개로, 기계가 검증할 수 있는 신호를 요구한다.

이것은 계산 가능한 "위험"을 다루던 언어의 연장선이다. 검증할 수 없는 표시는 집행할 수 없고, 집행할 수 없는 의무는 선언에 그친다.

격차는 실무에서 곧바로 돈으로 환산된다. 워터마킹이나 콘텐츠 출처 표준을 구현하려면 상당한 엔지니어링 투자가 든다. 안내 문구 삽입은 화면에 한 줄 넣는 일이다.

국내 시장만 본다면 이 격차를 못 느낄 수 있다. 그러나 사용자가 그 이미지를 캡처해 해외 플랫폼에 올리는 순간, 기계판독 신호가 없는 콘텐츠는 EU 기준에서 추적 불가능한 익명의 합성물이 된다. 국경을 넘는 것은 회사가 아니라 콘텐츠다.

한국에 없는 두 항목, Art.50(3)의 감정인식 고지와 Art.50(5)의 공익 텍스트 고지도 같은 결을 갖는다. 둘 다 당사자가 모르는 사이에 영향을 받는 상황을 겨냥한다. 사람이 스스로 알아차릴 수 없는 곳까지 법이 따라가 밝히라고 요구하는 일 — 이 역시 진흥법의 자연스러운 사정거리 밖에 있다.

밝힌다는 것의 두 얼굴

같은 동사, "밝히다"를 두 법이 함께 쓴다.

하나는 그 동사를 사람 앞에서 완성한다. 다른 하나는 기계 앞에서도 완성해야 끝났다고 본다.

다만 이 문법에는 두 법 모두에게 공통된 빈칸이 하나 있다.

고지와 표시는 "속지 않을 권리"를 지켜 주지만, 대화의 내용이 좋은지는 묻지 않는다.

밤 열한 시의 그 대화가 위로였는지 독이었는지, 어느 표시도 말해 주지 않는다. AI라고 밝힌 챗봇이 해로운 조언을 해도 표시 의무는 충족된 것이다. 투명성은 판단의 재료를 넘겨줄 뿐, 판단 자체를 대신하지 않는다.

그런데 투명성에는 또 다른 층이 있다. 스스로 밝히는 것과 남이 확인해주는 것은 다르다. 다음 장에서 "자율"이라는 이름을 달고 있지만 실제로는 규제로 작동하는 장치를 연다.

실무자를 위한 체크박스

- ☐ 생성형 기능이 있다면 출력물에 **기계판독 신호**(워터마크·메타데이터·서명)를 심을 계획이 있는가. 화면 안내 문구만으로는 EU 요건을 채우지 못한다
- ☐ 2026년 8월 2일 이후 출시분에는 경과조치가 없다. 신규 릴리스 일정과 워터마킹 구현 일정을 맞춰 두었는가
- ☐ 감정인식이나 생체분류 기능이 부수적으로라도 들어 있다면, 한국에는 없지만 EU에는 있는 고지 의무(Art.50(3))를 확인했는가
- ☐ 우리가 모델 제공자인가 앱 제공자인가. 제공자라면 기술문서·저작권 정책·학습데이터 요약 공개 의무가 2025년 8월부터 이미 적용 중이다
- ☐ 남의 모델을 받아 쓴다면, 그쪽에서 받는 정보가 우리 기술문서를 채우기에 충분한지 계약서에서 확인했는가

주(註)

1. AI Act Article 50(1)-(5).
2. 인공지능기본법 제31조제2항 및 시행령 제23조제2항. 기계판독 가능한 방법과 안내 문구 병행 중 선택이 가능하다.
3. AI Act Article 50(2). 워터마크·메타데이터·디지털서명 등 콘텐츠 출처 표준(C2PA) 계열 기술이 이행 수단으로 논의된다.
4. Digital Omnibus의 Article 111(2) 신설 조항. 2026년 8월 2일 이전에 이미 시장에 출시된 합성콘텐츠 생성 AI에만 2026년 12월 2일까지 경과조치를 준다. 신규 출시 시스템은 유예가

없다. 이사회 공식 채택과 관보 게재 전 단계이므로 집필 시점 재확인 필요하다. [검토 필요]

5. AI Act Article 53(범용 AI 모델 제공자의 기술문서·다운스트림 정보·저작권 정책·학습데이터 요약 의무), 적용일 2025년 8월 2일.
6. AI Act Article 5(1)(a)(b) 및 관련 전문. 조작이나 취약성 악용으로 인한 중대한 피해에 "신체적, 심리적 건강 또는 금전적 이익"에 대한 중대한 부정적 영향이 명시돼 있다.
7. Directive (EU) 2024/2853(개정 제조물책임지침) Article 4(1)에서 소프트웨어를 제품으로 정의하고, Article 6에서 의학적으로 인정된 정신건강 손상을 손해에 포함하며, Article 10(4)에서 기술적 복잡성으로 입증에 지나치게 어려운 경우 법원의 재량적 추정을 허용한다.

7장. 자율의 얼굴을 한 규제 — 검·인증과 우선구매

이 장을 한 문장으로

한국법은 검·인증을 "민간 자율"이라 부른다. 그러나 그것을 공공조달에 묶는 순간, 자율은 지렛대를 얻는다.

시장에 문이 하나 있다고 하자. 이 문을 여는 방법은 두 가지다.

하나는 **법이 잠그는 문**이다. 열쇠가 없으면 문 자체가 열리지 않는다. 예외가 없다.

다른 하나는 **돈이 여는 문**이다. 열쇠가 없어도 문은 열린다. 다만 문 안쪽의 특정 구역에는 들어갈 수 없다.

둘 다 "열쇠 없이는 어렵다"는 압력을 준다. 그러나 압력이 걸리는 범위가 다르다.

EU는 앞의 문을 만들었고, 한국은 뒤의 문을 만들었다.

다시 펼치는 그 문장

1장에서 한 문장을 인용했다. 한국 인공지능기본법의 제정이유서가 검·인증 조항을 설명하는 문장이었다.

"...민간이 자율적으로 추진하는 인공지능 안전성·신뢰성 검·인증 및 영향평가에 대한 지원을 하고..."^[1]

그때는 이 문장을 "법이 스스로 규제 조항을 진흥의 언어로 번역한 사례"로 읽고 지나갔다.

이번에는 그 문장을 끝까지 따라가 본다. 정말로 자율이기만 한가.

세 조문이 만드는 하나의 경로

한국 쪽 장치는 조문 하나가 아니다. 세 조문이 이어진 경로다.

1. 제30조 검·인증 지원 — 국가는 민간 자율 검·인증 활동을 지원함(제1항). 고영향 AI를 제공할 때 사전 검·인증은 노력 의무임(제3항). 그런데 제4항이 갈림길임 — 국가기관등이 고영향 AI를 이용할 때는 검·인증 받은 것을 우선 고려해야 함.^[2]
2. 시행령 제15조④⑤ 국가기관등 우선구매 — 법 제16조제3항의 위임을 받아, 과기정통부장관이 확인한 제품·서비스를 우선구매 대상으로 정함. 확인 기준과 절차는 행정안전부와 협의해 별도 고시로 정하는데, 조문 신설은 2026년 7월 20일이고 고시는 시행령 시행일과 같은 2026년 7월 21일에 나왔음 — 「인공지능제품·서비스 확인 절차 운영에 관한 고시」(과학기술정보통신부고시 제2026-50호).^[3]
3. 제35조 영향평가 — 역시 노력 의무이지만, 국가기관등은 영향평가를 실시한 제품·서비스를 우선 고려해야 함.^[4]

세 조문을 한 문장으로 이으면 이렇다.

"검·인증도 영향평가도 법적으로는 안 해도 그만이다. 그러나 안 하면 정부에는 못 판다."

왜 하필 이 둘만 노력 의무인가

1장에서 정리한 사실을 여기서 다시 꺼내야 한다.

같은 법의 제31조(투명성)·제32조(프론티어 안전성)·제34조(고영향 사업자 6책무)·제36조(국내대리인)는 모두 "이행하여야 한다"는 의무다. 어기면 사실조사와 시정명령이 오고 과태료가 붙는다. ^[5]

노력 의무로 남은 것은 제30조제3항 검·인증과 제35조 영향평가, 딱 둘이다.

둘 다 시장에 내놓기 전에 제3자가 확인하는 절차다.

한국법은 의무 전반을 피하지 않았다. 사전 관문 하나를 피했다.

그리고 그 자리를 조달이라는 다른 장치로 메웠다.

EU는 문 전체를 잠근다

EU의 대응 장치는 구조가 다르다.

Article 40은 조화표준을 지키면 요건 준수를 추정해 준다. Article 43은 고위험 AI에 강제 적합성평가를 요구한다. 통과하지 못하면 CE 마킹을 달 수 없고, CE 마킹이 없으면 EU 역내 시장 어디에도 낼 수 없다. ^[6]

공공기관에 팔든 개인 소비자에게 팔든 마찬가지다. 국경을 넘는 순간 같은 문을 통과해야 한다. 시장 전체가 하나의 문이다.

한국의 문은 다르다. 검·인증이 없어도 민간 시장에는 자유롭게 팔 수 있다. 문이 잠기는 곳은 공공조달이라는 한 구역뿐이다.

지렛대의 크기

한 구역뿐이라고 해서 이 장치가 약하다고 단정할 수는 없다. **정부는 작은 구매자가 아니다.**

특히 지금 한국은 정부의 AI 전환이 정책의 실질 동력으로 떠오르는 국면이다. 시행령 개정에서 행정안전부가 학습용데이터·기술도입지원·우선구매 확인 고시에 **공동 주체로 새로 등장한 것이 그 신호다.**^[7]

공공부문이 AI를 대규모로 도입하려는 참이라면, "정부에 팔 자격"은 결코 작은 시장이 아니다. 오히려 검·인증 없이는 지금 가장 빠르게 커지는 시장에서 통째로 배제된다는 뜻이 된다.

법적으로는 자율, 시장 구조상으로는 사실상의 관문. 이것이 이 장 제목의 의미다.

이 지렛대가 실제로 어떻게 작동할지는 한 장면으로 그려 볼 수 있다.

어느 부처가 민원 상담에 쓸 AI를 도입한다고 하자. 조달 공고에 "검·인증을 받은 제품을 우선 고려한다"는 한 줄이 붙는다. 법적으로는 받지 않은 제품도 응찰할 수 있다. 그러나 평가표에 가점이 있고 심사위원이 그 항목을 본다면, 받지 않은 회사는 사실상 들러리가 된다.

그다음이 진짜다. 한 부처가 그렇게 하면 다른 부처가 따라 하고, 지방자치단체가 따라 하고, 공공기관이 따라 한다. **조달 관행은 법보다 빠르게 복제된다.** 그 시점이 되면 "노력 의무"라는 법문은 현실을 설명하지 못한다.

그리고 이 확산에는 아무 입법 절차도 필요하지 않다. 고시 하나와 평가표 한 줄이면 충분하다.

[필자 분석] 강제하지 않고 유도하는 설계

이 설계를 "무늬만 자율"이라고 폄하하는 것은 정확하지 않다. **의도적으로 설계된 다른 종류의 정책 도구**로 읽는 편이 맞다.

4장에서 확인했다. 법으로 강제하면 산업 육성법이라는 정체성이 흔들린다. 그런데 아무 지렛대도 없으면 검·인증 참여율은 낮을 수밖에 없다.

그래서 한국이 택한 길은 **법적 강제 대신 시장의 지렛대**다. 국가가 스스로 구매자가 되어 구매 조건으로 검·인증을 요구한다. 강제가 아니라 유인이다. 진흥법의 문법 안에서도 작동하는 규제 효과를 설계한 셈이며, 4장에서 본 "금지는 진흥법의 문법이 아니다"라는 진단에 대한 한국식 우회 해법이기도 하다.

이유가 하나 더 있다. 그 단어 자체에 답이 들어 있다.

법은 이 활동을 "검증·인증 활동"이라 부르고 그것을 "검·인증등"으로 줄였다. 인증과 검증을 나란히 놓은 것이다.

무심코 지나치기 쉬운 조합이지만, 둘은 같은 것이 아니다.

인증은 이미 정해진 규격이 잣대다. "이 규격에 맞는가"를 묻는다. 묻는 쪽과 답하는 쪽이 같은 자를 들고 있고, 그래서 제3자가 판정할 수 있다. 통과와 불통과가 갈린다.

검증은 그 규격이 없을 때 쓴다. 잣대가 없으니 만든 쪽이 내건 주장을 잣대로 삼는다. "당신 말이 맞는가"를 묻는 것이다. 주장이 달라지면 검증의 내용도 달라진다. 회사마다 다른 자를 들고 오는 셈이다.

이 차이가 왜 결정적인가.

0-1에서 본 40년 설계의 핵심 장치가 조화표준이었다. 법은 "안전해야 한다"는 목표만 쓰고, 표준이 "이렇게 만들면 안전한 것으로 본다"는 구체적

방법을 제공한다. 표준이 잣대를 먼저 깔아 두기 때문에 인증이 성립하고, 인증이 성립하기 때문에 문을 잠글 수 있다.

EU에서 verification은 인증과 나란한 별개 활동이 아니다. 적합성평가 모듈의 이름이다. 잣대가 이미 있으니 verification은 그 잣대를 대는 방법 중 하나로 인증 절차 안에 들어가 있다.

한국에서는 검증이 인증과 병렬로 놓였다. 나란히 놓였다는 것은 아직 잣대가 없어서 둘 중 무엇을 쓸지 정하지 못했다는 뜻이다.

잣대가 없으면 무엇을 강제할지부터 특정할 수 없다. 무엇을 어겼는지 정의할 수 없는 의무는 집행할 수 없고, 집행할 수 없는 의무는 법정에서 버티지 못한다. 그래서 검증을 함께 묶었고, 강제 대신 유인을 택했다.

검·인증이라는 합성어는 편의상의 축약이 아니다. 규격이 아직 없다는 상태를 법이 문장 구조로 기록해 둔 것이다.

거꾸로 말하면 이렇다. 한국이 언젠가 EU식으로 문을 잠그려 한다면, 손봐야 할 것은 제30조의 "노력하여야 한다"가 아니다. 그 아래에 깔릴 표준이다. 조문의 서술어를 바꾸는 일은 하루면 되지만 잣대를 만드는 일은 몇 년이 걸린다. 2장에서 본 "기초 없이 새 층만 올린 형국"이 여기서 구체적인 모습으로 드러난다.

이것은 한국만의 발명이 아니다

여기서 짚어 둘 것이 있다. "자율의 얼굴을 한 규제"를 한국의 특이한 우회로로만 보면 절반만 읽은 것이다. EU도 같은 기법을 쓴다.

범용 AI 모델에 적용되는 실행규범(Code of Practice)이 그것이다. 형식은 자율 서명이고 법이 강제하는 문서가 아니다. 그런데 서명하지 않은 공급자는 관련 의무를 각자 개별적으로 입증해야 하고, 서명한 공급자는 그 부담 없이 준수로 추정받는다.

법적으로는 자율, 실질적으로는 경로를 하나로 몰아가는 지렛대. 문법이 같다.

다만 지렛대의 위치가 다르다. 한국은 공공조달이라는 시장의 힘을 쓰고, EU는 입증 부담의 비대칭을 쓴다. 한쪽은 돈으로 밀고 다른 쪽은 절차 비용으로 민다. 8장에서 이 대비를 다시 만나게 된다.

두 가지 한계

그러나 한국식 설계에는 EU식 문과 다른 근본적 한계가 있다.

첫째, 사정거리가 공공조달에서 멈춘다. 순수 민간시장을 겨냥한 AI, 해외 소비자를 겨냥한 AI, 정부와 거래가 없는 스타트업의 AI는 이 지렛대의 사정권 밖에 있다. EU의 문은 시장 전체가 지나가야 하는 관문이지만, 한국의 문은 정부라는 한 구역의 검문소일 뿐이다.

둘째, 지렛대의 무게는 예상과 다른 축에서 정해졌다. 이 원고를 처음 쓸 때만 해도 우선구매의 실제 기준을 정할 고시가 없었다. 그 고시는 시행령과 같은 날 나왔다.^[3:1] "기준이 느슨하면 상징에 그치고, 촘촘하면 국가 인증제에 가까워진다"고 앞서 물었던 갈림길은, 조문을 읽어 보니 둘 중 하나로 떨어지지 않는다. 촘촘하되, 안전성이 아닌 다른 것을 촘촘히 쥔다.

고시가 확인기관에 요구하는 심사 기준은 일곱 가지다. 요약하면 전부 한 질문으로 수렴한다 — 이 제품에 정말 AI 연산체계가 들어 있고, 그것이 실제로 작동하는가. 학습·추론의 특성을 갖췄는가, 입력이 그 체계를 거쳐 의미 있는 출력을 내는가, 구조도상 위치가 명확한가, 정상 동작하는가. 정확도나 공정성, 안전성, 견고성을 묻는 기준은 하나도 없다.

즉 이 관문이 막으려는 것은 위험한 AI가 아니라 가짜 AI다. "우리 제품은 AI 기반입니다"라고 말해 놓고 실제로는 규칙 기반 자동화에 불과한 경우, 조달 우대를 받지 못하도록 걸러내는 장치에 가깝다. 확인은 부처가 직접

하지 않는다. 확인기관(한국인공지능진흥협회)이 신청을 받고, 심사기관(한국정보통신기술협회)이 기술 심사를 맡는 이원 구조다. 국가가 기관을 지정해 심사를 위임하는 형태는 0-2에서 본 유럽의 검증기관(Notified Body) 배치와 겉모습이 닮았다. 그러나 유럽의 검증기관은 **안전한가**를 묻고, 이 확인기관은 **진짜인가**를 묻는다. 같은 형태의 장치가 다른 문에 서서, 다른 질문을 던지는 것이다.

조문은 확정됐고 고시도 나왔다. 그 고시가 답한 것은 "얼마나 안전해야 우대받는가"가 아니라 "무엇을 AI라고 부를 것인가"였다. 0-3에서 선언한 독법이 여기서 그대로 적용된다.

정부가 채우지 않은 그 빈자리를 민간이 채우려는 움직임도 시작됐다. 산·학·연 협의체가 통합 검·인증 체계를 준비 중이고, 2026년 7월부터 시범 인증에 들어갔다.^[8] 국가가 강제하지 않은 자리에서, 시장이 스스로 기준을 만들려는 참이다.

두 개의 문, 두 개의 사정거리

법이 잠그는 문과 돈이 여는 문. 어느 쪽이 더 규제에 가까운지를 묻는 것은 잘못된 질문일 수 있다.

진짜 질문은 이것이다. **그 문이 얼마나 넓은 시장을 지키고 있는가.**

그리고 이 조달의 지렛대는 결국 다른 질문과 만난다. 얼마나 큰 AI를 얼마나 무겁게 다룰 것인가. 다음 장에서 규모 그 자체가 문턱이 되는 지점을 연다.

실무자를 위한 체크박스

- ☐ 공공 매출 비중이 얼마인가. 그 비중이 검·인증 투자의 손익분기를 결정한다
- ☐ 우선구매 확인 고시가 나오면 **즉시 검토할 담당자를 지정해 두었는가**. 기준이 나온 뒤 준비를 시작하면 이미 늦다 (14장 모니터링 항목)
- ☐ 지금 받을 수 있는 검·인증이 무엇인지 파악했는가. 잣대가 확정되기 전이라 선택지 자체가 유동적이다
- ☐ EU 진출 계획이 있다면, 조화표준 기반 적합성평가와 국내 검·인증 중 **어느 쪽 문서가 다른 쪽에 재사용되는지** 매핑했는가
- ☐ 범용 AI 모델을 제공한다면 EU 실행규범 서명 여부를 결정했는가. 서명하지 않으면 개별 인증 부담이 우리에게 온다
- ☐ 대외 메시지를 만들 때 "민간 자율이니 규제가 아니다"라고 설명하고 있지 않은가. 조달 담당 부처는 그렇게 보지 않는다

주(註)

1. 인공지능기본법 원 제정(법률 제20676호) 제정이유 주요내용 (자)항.
2. 인공지능기본법 제30조 제1항·제3항·제4항. "검·인증등"은 제1항이 만든 법률상 정의어다 — "자율적으로 추진하는 검증·인증 활동(이하 '검·인증등'이라 한다)". 한국어 "검증"에는 영어 두 단어가 겹쳐 있다. verification은 명세된 요구사항을 충족했는지("제대로 만들었는가")를, validation은 의도한 사용 목적을 충족했는지("맞는 것을 만들었는가")를 묻는다. EU에서 verification은 인증과 나란한 별개 활동이 아니라 적합성평가 모듈의 이름이다(Module F-F1 product verification, Module G unit verification).
3. 인공지능기본법 시행령 제15조④제1호·⑤(2026년 7월 20일 신설). 법 제16조제3항의 위임을 받은 국가기관등 우선구매 대상은 과기정통부장관 확인 제품·서비스이며, 확인 기준·절차는 행정안전부 협의 고시로 위임됐다. 위임을 받아 「인공지능제품·서비스 확인 절차 운영에 관한 고시」(과학기술정보통신부고시 제2026-50호, 2026-07-21 제정·시행)가 원문으로 확인됐다(raw/docs 인제스트, 2026-08-05). 확인기관은 한국인공지능진흥협회(법 제26조①), 심사기관은 한국정보통신기술협회(방송통신발전 기본법 제34조). 확인기준 7개(제5조)는 전부 "인공지능처리 연산체계가 실제로 적용·작동하는가"를 묻는 기술적 존재 확인이며, 정확도·공정성·안전성·견고성 등 성능·품질 기준은 포함되지 않는다. 처리기간은 접수 후 15일(제품·서비스 복잡성 사유로 30일 범위 내 1회 연장 가능). 신청 전 조달청

목록정보시스템에서 물품목록번호를 사전 확보해야 한다(제4조②). 거짓·부정 확인은 취소 사유다(제8조). 상세는 인공지능제품서비스확인절차고시 참조.

4. 인공지능기본법 제35조. 영향평가는 노력 의무이며, 국가기관등은 영향평가를 실시한 제품·서비스를 우선 고려한다.
5. 의무/노력의무 스펙트럼은 1장 주석 7 참조. 위반 시 제40조 사실조사·중지·시정명령과 제43조 과태료(상한 3천만원)로 연결된다.
6. AI Act Article 40(조화표준 준수추정) 및 Article 43(적합성평가 절차, 미이행 시 CE 마킹·시장출시 불가).
7. 인공지능기본법 시행령 제13조③·제15조②③⑤(행정안전부 공동 주체 신설, 2026년 7월 20일).
8. 서울경제, "제각각 AI 안전성 검·인증...연내 통합체계 나온다", 2026-07-28. 산·학·연 협의체 "AI 신뢰성 얼라이언스"가 2026년 7월부터 6개 기업 대상 시범인증에 착수, 2026년 12월 본격 시행 목표. 기존 국내 검·인증 제도는 TTA·KTL·KSA 등 7개 이상 난립. NIA 채은선, "인공지능기본법의 주요 내용" 세미나(AI 제품안전 혁신 TF 제6차 전체 회의, 2026-08-06) 인용. 상세는 검증과-인증 § 6 참조.

8장. 가장 큰 모델의 문턱 — 10^{26} 과 10^{25} 사이

이 장을 한 문장으로

두 법이 "이 AI는 특별히 감시해야 한다"고 선을 긋는 지점이 자릿수 하나만큼 어긋나 있고, 그 하나가 열 배다.

10^{25} 와 10^{26} . 자릿수 하나 차이다.

눈으로 보면 사소해 보인다. 25와 26. 하나 늘었을 뿐이다.

그런데 지수는 그렇게 작동하지 않는다. 10^{25} 는 1 뒤에 0이 스물다섯 개, 10^{26} 은 스물여섯 개다. 후자는 전자의 정확히 열 배다. 100원과 1,000원의 차이가 표기상 숫자 하나로 줄어드는 것과 같다.

두 법이 가장 큰 AI를 가르는 선이 바로 이 자릿수 하나의 간극 위에 있다.

왜 하필 연산량인가

먼저 단위부터 풀어야 한다.

FLOPs는 부동소수점 연산 횟수를 뜻한다. 컴퓨터가 소수점이 붙은 숫자를 더하고 곱한 총 횟수다. AI 모델을 훈련한다는 것은 결국 이 계산을 천문학적으로 반복하는 일이고, 그 누적 횟수가 훈련에 들어간 자원의 크기를 대략 말해 준다.

여기서 물어야 할 것이 있다. 왜 두 법 모두 하필 연산량을 기준으로 삼았는가.

위험한 것은 연산량이 아니다. 모델이 무엇을 할 수 있는가, 즉 능력이다. 그런데 능력은 직접 재기 어렵다. 어떤 시험을 몇 점 받으면 위험한 모델인지에 합의된 잣대가 없고, 그 시험을 만드는 일 자체가 연구 주제다.

반면 연산량은 세기 쉽다. 훈련에 쓴 칩의 수와 시간을 곱하면 나온다. 개발사가 스스로 알고, 회계 장부에도 흔적이 남고, 사후에 검증할 여지도 있다.

즉 연산량 문턱은 위험을 재는 자가 아니라 위험을 대신 가리키는 표지다. 규제가 직접 재고 싶은 것을 잴 수 없을 때 쓰는 대리 지표다.

이 선택에는 대가가 따른다. 대리 지표는 언제나 원래 재려던 것과 조금씩 어긋난다. 적은 연산으로 훈련한 위험한 모델이 있을 수 있고, 막대한 연산을 쓰고도 무해한 모델이 있을 수 있다. 두 법 모두 이 어긋남을 알고 있었고, 각자 다른 방식으로 보완했다. EU는 집행위원회가 직권으로 추가 지정할 수 있는 길을 열어 두었고,^[1] 한국은 연산량 외에 최첨단 기술 해당 여부와 위험의 광범위성을 요건으로 함께 걸었다.

문턱은 자동판정기가 아니라 1차 거름망이다. 이 점을 기억해 두면 뒤의 이야기가 다르게 읽힌다.

같은 모델, 다른 신분

한국 인공지능기본법 제32조는 누적 연산량이 대통령령 기준 이상인 AI 시스템에 별도 의무를 부과한다. 시행령이 정한 그 기준이 10^{26} FLOPs다.^[2]

EU AI Act Article 51은 10^{25} FLOPs를 시스템적 위험의 추정 기준으로 삼는다.^[1:1]

열 배의 간극. 그리고 이 간극 사이에 실제로 걸리는 모델들이 있다.

EU 기준은 넘지만 한국 기준에는 못 미치는 구간이 존재한다는 뜻이다. 같은 모델이 국경을 넘는 순간 신분이 바뀐다. 유럽에서는 특별 감시 대상, 한국에서는 평범한 AI.^[3]

물론 최상위 모델은 두 기준을 모두 넘는다. 문턱의 차이가 실제로 체감되는 곳은 정확히 두 자릿수 사이에 낀 구간이다.

문턱을 넘은 뒤 무엇을 하는가

문턱의 위치만 다른 것이 아니다. 넘은 뒤 요구받는 것의 구체성도 다르다.

한국 제32조가 요구하는 것은 둘이다.^[4]

1. 수명주기 전반에 걸친 위험 식별·평가·완화
2. 안전사고 모니터링·대응 위험관리체계 구축

그리고 과학기술정보통신부에 이행 결과를 제출한다. 구체적인 이행 방식은 고시로 정하는데, 무엇을 어떻게 평가해야 하는지의 세부는 아직 확정 전이다.

EU Article 51에서 55까지는 항목을 훨씬 촘촘하게 나열한다.^[5]

1. 시스템적 위험 평가 — 모델 자체의 위험 특성 분석임
2. 적대적 테스트 — 실제로 공격을 가해 보고 취약점을 찾는 절차임
3. 중대 사고 보고 — 정해진 시한 안에 AI Office와 회원국 당국에 보고함
4. 사이버보안 요건
5. 실행규범(Code of Practice) — 위 의무들의 표준 이행 경로임

한국의 두 항목은 목표를 제시하는 **선언형**이다. EU의 목록은 절차를 요구하는 **체크리스트형**이다.

3장에서 확인한 "영향 대 위험"의 구도가 가장 큰 모델 앞에서도 똑같이 반복된다.

차이가 가장 선명한 것은 두 번째 항목, **적대적 테스트**다.

일반적인 품질 검사는 제품이 제대로 작동하는지를 본다. 적대적 테스트는 반대다. **제품을 망가뜨리려고 작성하고 덤빈다.** 금지된 답을 끌어내려 프롬프트를 비틀고, 안전장치를 우회할 경로를 찾고, 모델이 하지 말아야 할 일을 하게 만들려 시도한다. 그 시도가 성공하면 그것이 취약점이다.

이 절차에는 별도의 인력과 시간이 든다. 만든 사람이 자기 모델을 공격하는 일은 심리적으로도 조직적으로도 쉽지 않아서, 대개 독립된 팀을 둔다. **한국 제32조에는 이에 해당하는 항목이 없다.** "위험을 식별·평가·완화하라"는 문장이 적대적 테스트를 포함한다고 읽을 수는 있지만, 읽어야만 나온다.

시스템적 위험이라는 개념도 짚어 둘 필요가 있다. 개별 사용자가 입는 피해가 아니라, 모델 하나가 사회 전체 규모로 일으킬 수 있는 파급을 가리킨다. 하나의 기반 모델 위에 수천 개의 서비스가 올라가는 구조에서는 그 모델의 결함이 한 지점의 사고로 끝나지 않는다. 가장 큰 모델을 따로 떼어 규율하는 이유가 여기 있다.

자율이라는 이름의 재등장

다섯 번째 항목인 실행규범을 다시 볼 필요가 있다. 7장에서 이미 만난 그 구조다.

형식은 **자율 서명**이지만, 서명하지 않은 공급자는 같은 의무를 혼자 입증해야 한다. 자율과 강제 사이에 입증 부담이 놓여 있는 구조다.

한국은 공공조달이라는 시장의 힘을 지렛대로 쓰고, EU는 입증 부담의 비대칭을 지렛대로 쓴다.

같은 문법, 다른 지렛목이다. 자율의 얼굴을 한 규제는 한국만의 발명품이 아니었다.

다만 실무적 함의는 정반대에 가깝다. 한국의 지렛대는 **정부에 팔지 않으면 피할 수 있다**. EU의 지렛대는 피할 방법이 없다 — 서명하지 않는 선택지는 "의무가 없다"가 아니라 "혼자 다 입증하라"다.

[필자 분석] 관대함은 고정되지 않는다

10^{26} 이라는 숫자를 "한국이 산업을 보호하려고 정한 여유로운 기준"으로 읽기 쉽다. 틀린 해석은 아니다. 그러나 이 읽기에는 **시간이라는 변수가 빠져 있다**.

연산량 기준은 정적인 숫자다. 그런데 컴퓨팅 비용은 정적이지 않다. 하드웨어가 발전할수록 같은 돈으로 더 많은 연산을 살 수 있다.

오늘 10^{26} FLOPs를 훈련하려면 막대한 자본이 필요하지만, 몇 년 뒤에는 훨씬 적은 자본으로 같은 연산량에 도달한다.

문턱의 절대 위치는 고정돼 있어도, 그 문턱에 도달하는 데 드는 진입장벽은 계속 낮아진다. 오늘의 관대한 기준이 몇 년 뒤에는 웬만한 기업도 넘는 기준이 될 수 있다는 뜻이다.

한국 시행령에 이 문제를 의식한 흔적이 있다. 프론티어 기준을 포함한 규제 전반을 **3년마다 재검토**하도록 정해 두었다.^[6] 정적인 숫자를 동적으로 관리하겠다는 장치다.

문제는 재검토의 결과가 어느 방향으로 갈지 아무도 모른다는 것이다. 기준을 EU 쪽으로 좁힐 수도, 산업 보호 논리로 그대로 둘 수도 있다.

그래서 이 장에서 가장 중요한 교훈은 "지금 몇 배 차이인가"가 아니다. 이 숫자가 시간이 지나도 같은 의미를 가질 것인가다.

3장에서 본 언어의 차이, 7장에서 본 자율의 지렛대 위에 프론티어 규제는 층을 하나 더 얹는다. 숫자로 그은 선은 안정적으로 보이지만, 그 선이 가리키는 실제 위치는 기술이 움직이는 만큼 함께 움직인다.

자릿수 하나가 가리키는 것

10^{25} 와 10^{26} . 자릿수 하나의 차이가 오늘은 몇 개의 모델을 가르는 선이지만, 내일은 훨씬 많은 모델을 가르는 선이 될 수 있다.

정적인 법이 동적인 기술을 따라잡으려면, 문턱 자체보다 문턱을 다시 그을 수 있는 장치가 더 중요하다.

이제 남은 것은 하나다. 문턱을 넘었는데 의무를 어겼다면, 무슨 일이 벌어지는가.

실무자를 위한 체크박스

- ☐ 우리 모델의 누적 훈련 연산량을 **계산해 두었는가**. 계산 방식 자체가 아직 고시 대기 중이므로, 산정 근거를 남겨 두는 것이 먼저다
- ☐ 두 문턱 사이 구간에 있다면, **EU에서만 시스템적 위험 대상**이 된다. 유럽 진출 일정에 이 의무를 반영했는가
- ☐ 적대적 테스트를 수행할 인력과 절차가 사내에 있는가. 한국에는 없지만 EU에는 명시된 항목이다
- ☐ 실행규범 서명 여부를 결정했는가. 서명하지 않으면 개별 입증 부담이 우리에게 온다
- ☐ 3년 주기 재검토로 한국 기준이 낮아질 가능성을 사업계획의 리스크로 잡아 두었는가

주(註)

1. AI Act Article 51(2). 훈련 연산량 10^{25} FLOPs 초과 시 시스템적 위험으로 추정한다. 집행위원회 직권이나 과학패널 경보를 통한 추가 지정도 가능하다.
2. 인공지능기본법 제32조(인공지능 안전성 확보 의무) 및 시행령 제24조. 시행령 제24조①은 세 요건을 모두 충족할 것을 요구한다 — 학습에 사용된 누적 연산량 10^{26} 부동소수점연산 이상, 최첨단 인공지능기술 적용, 사람의 생명·신체 안전 및 기본권에 광범위하고 중대한 영향을 미칠 우려. 연산량은 세 요건 중 하나일 뿐 단독 판정 기준이 아니다(법제처 국가법령정보 조문 대조, 2026-08-05).
3. 개별 모델의 정확한 훈련 연산량은 개발사가 공식 공개하지 않는 경우가 많다. 여기서는 두 문턱 사이에 걸리는 모델군이 존재한다는 구조적 사실만 언급하며, 특정 모델의 수치는 인용하지 않았다. [검토 필요]
4. 인공지능기본법 제32조제1항 제1호·제2호(수명주기 전반의 위험 식별·평가·완화 / 안전사고 모니터링·대응 위험관리체계 구축) 및 제2항(이행 결과 제출). 제3항은 "과학기술정보통신부장관은 …구체적인 이행 방식 및 …결과 제출 등에 필요한 사항을 정하여 고시하여야 한다"고 규정한다 — 재량이 아니라 의무다. 그 고시는 2026-08-05 초회 기준 제정되지 않았다. 별도로 연산량 산정 방식도 시행령 제24조②가 고시에 위임했고 역시 미제정이다(부록 B).

5. AI Act Article 51(시스템적 위험 추정)·Article 55(추가 의무: 모델평가·적대적 테스트·사이버보안·중대사고 보고)·Article 56(실행규범).
6. 인공지능기본법 시행령 제33조(규제의 재검토). 2026년 1월 1일 기준 3년 주기이며 프론티어 연산량 기준이 대상에 포함된다.

9장. 어기면 어떻게 되는가 — 3천만원 vs 매출 7%

이 장을 한 문장으로

제재의 크기 차이는 숫자의 문제가 아니라 태생의 문제이며, 먼저 시행되는 법이 반드시 더 준비된 법은 아니다.

법의 진심은 벌칙 조항에서 드러난다.

앞의 여덟 장에서 본 모든 차이 — 분류 방식, 금지 목록의 유무, 표시의 수신자, 자율의 실질, 문턱의 위치 — 는 결국 하나로 수렴한다. **안 지켰을 때 얼마나 아픈가.**

여기서 격차가 가장 노골적으로 드러난다.

숫자로 보는 마지막 격차

한국 인공지능기본법 제43조와 시행령 별표 2는 과태료를 이렇게 정한다.^[1]

1. 투명성 불이행(제31조제1항) — 1차 500만원, 2차 1,000만원, 3차 이상 1,500만원임.
2. 국내대리인 미지정(제36조제1항) — 2,000만원 정액임.
3. 시정명령 불이행(제40조제3항) — 1차 1,000만원, 2차 2,000만원, 3차 이상 3,000만원임.
4. 금지 AI 위반 — 해당 조항 자체가 없음(4장).

5. 프론티어 모델 의무 위반 — 과태료 규정 자체가 없음(8장).

법률 전체를 통틀어 가장 무거운 처분이 3,000만원이다.

EU AI Act Article 99와 101은 이렇다.^[2]

1. 금지 AI 위반 — 3,500만 유로 또는 전 세계 연간 매출의 7퍼센트 중 더 큰 쪽임.
2. 고위험 AI 의무 위반(투명성 포함) — 1,500만 유로 또는 매출의 3퍼센트임.
3. 범용 AI 모델 의무 위반 — 1,500만 유로 또는 매출의 3퍼센트임.

두 법 다 위반 유형별로 등급을 나눈다는 형식은 같다. 그런데 한쪽은 정액이고 다른 쪽은 매출에 연동된다.

매출 연동은 기업이 클수록 처분도 함께 커진다는 뜻이다. 정액 상한은 반대다. 기업이 클수록 처분은 상대적으로 가벼워진다. 매출 수십조 원 기업에게 3,000만원은 규제가 아니라 반올림 오차에 가깝다.

여기에 더 중요한 사실이 있다. 한국에서 가장 무거운 처분이 걸리는 위반은 금지 AI도 프론티어 모델도 아니다. **시정명령을 따르지 않은 것이다.** 실제적 위반이 아니라 절차 불응이 최고 처분을 받는 구조다.

과태료와 과징금은 다른 물건이다

숫자를 비교하기 전에 성격부터 갈라야 한다.

과태료는 행정 질서를 어긴 데 대한 제재다. 신고를 안 했다가나 표시를 빠뜨렸다가나 하는 의무 불이행에 붙는다. 액수를 법이 미리 정해 두고, 위반 횟수에 따라 올라간다. 어긴 사람이 그 위반으로 얼마를 벌었는지는 계산에 들어가지 않는다.

EU의 과징금은 다르다. 매출에 연동한다는 것은 위반으로 얻은 이익을 회수하겠다는 발상에 가깝다. 규제를 어기고 얻는 이득이 제재보다 크면 어기는 쪽이 합리적 선택이 되는데, 매출 연동은 그 계산을 뒤집기 위한 설계다.

그래서 한국의 3,000만원과 EU의 매출 7퍼센트는 크기가 다른 같은 물건이 아니라, 애초에 다른 물건이다. 앞은 질서 위반에 대한 벌이고, 뒤는 위반을 수지 맞지 않게 만들려는 장치다.

그 앞에 절차가 있다

과태료가 곧바로 날아오는 것도 아니다. 한국법의 집행은 단계를 밟는다.

과학기술정보통신부장관이 사실조사를 한다. 사무소나 사업장에 출입해 장부와 자료를 조사할 수 있다. 조사 결과 위반이 확인되면 중지 또는 시정명령을 내린다. 그 명령을 따르지 않을 때 비로소 최고액 과태료가 붙는다.^[3]

감경과 가중도 있다. 중소기업·벤처·소상공인은 2분의 1까지 감경받을 수 있고, 중대한 피해가 발생했거나 6개월 이상 위반이 이어졌으면 상한 안에서 2분의 1까지 가중된다.

실무적으로 이 구조가 뜻하는 바는 분명하다. 첫 통지에 성실히 대응하면 최고 처분까지 가지 않는다. 반대로 당국 문서를 방치하면 실제적 위반보다 절차 불응으로 더 크게 물린다.

시간표가 만드는 역설

숫자만큼 흥미로운 것이 시간표다. 두 법의 주요 적용일을 나란히 놓아 보자.^[4]

	한국	EU (원안)	EU (조정안)
투명성 의무	2026.07.21	2026.08.02	변경 없음
고위험 AI (Annex III)	2026.07.21	2026.08.02	2027.12.02
고위험 AI (Annex I, 제품내장)	2026.07.21	2027.08.02	2028.08.02

원안 기준으로는 한국과 EU가 2주 차이로 거의 나란히 출발한다.

그런데 EU가 고위험 AI 적용을 미루면서 그림이 완전히 달라진다. 한국의 고위험 AI 사업자 책무가 이미 시행되고 있는 동안, EU의 고위험 AI 요건은 1년 4개월을 더 기다린다.

한국이 EU보다 먼저 시행한다.

언뜻 "한국이 더 빠르고 앞선 규제를 갖췄다"는 이야기로 들린다. 정말 그런가.

[필자 분석] 먼저 도착한 것이 자량은 아니다

EU가 고위험 AI 시행을 미룬 이유는 게을러서가 아니다.

매출 7퍼센트라는 처분을 실제로 집행하려면, 그 전에 조화표준을 마련하고 검증기관을 지정하고 시장감시 인프라를 갖춰야 한다. 잣대 없이 그런 처분을 내리면 법정에서 버티지 못한다.

무거운 제재를 걸려면, 그 제재를 정당하게 발동할 수 있는 두꺼운 기반부터 완성해야 한다. EU가 늦게 도착하는 것은 그 두께를 쌓는 데 시간이 걸리기 때문이다.

한국이 먼저 도착한 이유도 게을러서가 아니다. 다만 이유가 다르다.

500만원에서 3,000만원 사이의 정액 과태료는 조화표준 체계나 대규모 시장감시 인프라 없이도 당장 집행할 수 있는 수준이다. 가벼운 제재는

무거운 기반을 요구하지 않는다.

얇은 법은 준비할 것이 적기 때문에 먼저 도착한다.

그러니 "한국이 EU보다 먼저 시행한다"는 사실은 두 가지로 동시에 읽힌다. 국내 기업 입장에서는 대응 부담이 EU보다 먼저 찾아온다는 실무적 사실이다. 그러나 규제 설계의 관점에서는 다른 신호다.

먼저 시행되는 법이 반드시 더 준비된 법은 아니다. 때로는 준비할 것이 적은 법이 먼저 도착할 뿐이다.

그리고 이 진단은 8장의 결론과 맞물린다. 프론티어 문턱을 3년마다 재검토하겠다는 장치가 있어도, 문턱을 넘은 자에게 물릴 제재가 없으면 그 재검토는 무엇을 위한 것인가. 한국법에는 프론티어 의무 위반에 대한 과태료 규정이 아예 없다. 문턱은 그었는데 문턱 너머에 아무것도 두지 않은 셈이다.

2부를 마치며

2부에서 우리는 돈과 시간이 실제로 드는 지점들을 봤다.

무엇을 표시해야 하고(6장), 무엇을 인증받아야 하고(7장), 어느 크기부터 특별 감시를 받고(8장), 어기면 얼마를 무는가(9장).

네 장을 관통하는 것은 하나다. 한국법의 조문은 EU와 비슷한 자리에 비슷한 이름으로 놓여 있지만, 그 조문 뒤에 서 있는 집행 장치의 무게가 다르다.

표시 의무는 있으나 기계판독을 요구하지 않고, 검·인증 조항은 있으나 조달 구역에서 멈추고, 프론티어 문턱은 있으나 제재가 없다. 조문이 없는 것이 아니라 조문 뒤가 비어 있다.

한줄 코멘트. 3천만원과 매출 7퍼센트의 격차는 숫자의 문제가 아니라 태생의 문제다. 그리고 누가 먼저 도착했는가보다 중요한 질문은, 도착했을 때 그 법이 실제로 작동할 준비가 되어 있었는가다.

두 법 모두 아직 완성된 이야기가 아니다. 한국은 진흥의 두께 위에 규제의 두께를 얼마나 더할 것인가. EU는 미뤄 둔 1년 4개월 뒤에 실제로 그 두꺼운 기반 위에서 매출 7퍼센트를 집행할 수 있을 것인가.

두 질문의 답은 아직 쓰이지 않았다.

그런데 답이 쓰이기 전에도 사고는 난다. 3부에서 잘못됐을 때 누가 물어주고 평소에는 누가 지켜보는지를 본다.

실무자를 위한 체크박스

- ☐ 우리 규모에서 3,000만원과 매출 3퍼센트 중 어느 쪽이 실질적 리스크인지 **환산해 보았는가**. 두 법의 체감 무게는 회사 크기에 따라 정반대로 뒤집힌다
- ☐ 한국에서 최고 처분이 걸리는 것은 **시정명령 불응**이다. 당국 문서에 대응할 창구와 기한 관리 절차가 있는가
- ☐ EU 과징금은 매출 연동이다. 모회사 연결 매출 기준인지 법인 단위인지 확인했는가
- ☐ 두 법의 시행일이 어긋난다는 사실을 제품 로드맵에 반영했는가. 한국 의무가 먼저 온다
- ☐ EU 고위험 요건이 미뤄진 기간을 **준비 기간으로 쓸 것인지 관망할 것인지** 결정했는가. 미뤄진 것은 시행일이지 요건이 아니다

주(註)

1. 인공지능기본법 제43조 및 시행령 별표 2(과태료 부과기준). 중소기업·벤처·소상공인 등은 2분의 1 감경이 가능하고, 중대 피해나 6개월 이상 위반은 상한 내에서 2분의 1 가중된다.
2. AI Act Article 99(3)(4) 및 Article 101. 금지 AI 위반은 3,500만 유로 또는 전 세계 연간 매출 7퍼센트 중 더 큰 금액이며, 그 아래 단계는 1,500만 유로 또는 3퍼센트다.
3. 인공지능기본법 제40조(사실조사 등). 과학기술정보통신부장관은 제31조제2항·제3항, 제32조제1항·제2항, 제34조제1항 위반과 관련해 조사할 수 있고, 사무소·사업장에 출입해 장부·서류·자료를 조사할 수 있으며, 중지 또는 시정을 명할 수 있다.
4. 인공지능기본법 제31조·제34조 시행일과 AI Act Article 113의 단계별 적용일. 조정안은 Digital Omnibus 반영 기준이며, 이사회 공식 채택과 관보 게재 전 단계이므로 집필 시점 재확인 필요하다. [검토 필요]

PART 3

3부 — 사고 이후: 책임과 감시

잘못됐을 때 누가 물어주고, 평소에는 누가 지켜보는가.

10장. AI가 잘못 판단했을 때 누가 책임지는가

이 장을 한 문장으로

AI의 불투명성은 만든 쪽의 구조적 이점이고, 두 법은 그 이점을 어디까지 깎을지에서 갈린다.

어느 기업이 AI 채용 시스템을 도입했다고 하자.

지원자 수천 명의 이력서를 AI가 1차로 분류한다. 70퍼센트는 자동 탈락하고, 30퍼센트만 사람이 검토한다.

A씨는 세 번 지원했고 세 번 모두 1차에서 떨어졌다. 경력이 나쁘지 않았다.

나중에 알려진 사실이 있다. 그 AI는 특정 대학 출신이나 특정 어휘 패턴에 과도한 가중치를 주고 있었다. 일종의 구조적 편향이다.

A씨가 할 수 있는 일은 무엇인가.

피해자의 세 가지 질문

이런 상황에서 피해자는 순서대로 세 가지를 묻는다.

1. 나는 왜 탈락했는가. — 설명을 요구할 권리가 있는가
2. AI가 나를 차별했다는 것을 어떻게 증명하는가. — 입증 책임이 누구에게 있는가
3. 누구를 상대로 무엇을 청구하는가. — 실제 구제 수단이 있는가

두 법의 답이 다르다. 그리고 이 장에서 가장 중요한 것은 EU조차 이 세 질문에 완전히 답하지 못했다는 사실이다.

첫 번째 질문 — 설명

EU AI Act Article 86은 고위험 AI의 결정에 영향을 받은 자연인에게 **명확하고 의미 있는 설명을 요구할 권리를 준다.**^[1]

핵심은 "의미 있는"이다. "AI가 판단했습니다"라는 한마디로는 안 된다. 어떤 요소가 어떻게 결정에 영향을 미쳤는지를 설명해야 한다. 사람이 다시 심사하도록 요청할 권리도 함께 붙는다.

한계도 분명하다. 고위험 AI에만 적용되고, 의무를 지는 것은 만든 쪽이 아니라 **그 AI를 실제로 쓴 쪽(배포자)**이다.

한국은 어떨까.

법 제31조제1항은 고영향·생성형 AI를 이용할 때 **AI에 기반해 운용된다는 사실을 사전에 고지하라고** 한다. 그런데 이것은 "AI가 관여한다"는 사실의 통보이지 결정의 **이유에 대한 설명**이 아니다.

결정 이유를 설명하라는 조항은 인공지능기본법에 없다.

개인정보보호법 제37조의2가 보조적으로 작동한다. 자동화된 결정에 대해 거부하고, 설명을 요구하고, 이의를 제기할 수 있다. 다만 이 조항은 **개인정보를 처리한 결정에만** 적용된다.^[2] 개인정보를 쓰지 않고 문서 텍스트만으로 판단하는 AI에는 미치지 않는다.

두 번째 질문 — 입증

설명을 받았다 해도 소송은 다른 문제다. "AI의 편향 때문에 내가 탈락했다"를 어떻게 증명하는가.

EU는 2022년에 이 문제를 정면으로 풀려 했다. **AI 책임지침**(AI Liability Directive) 제안이다.

설계된 장치는 이랬다. 기업이 AI Act 의무를 위반한 것이 확인되면, 법원은 그 위반과 피해 사이의 **인과관계를 추정**한다. 피해자가 AI 내부를 들여다볼 필요 없이, 기업이 자신은 의무를 제대로 지켰다고 반증해야 한다.

이 지침은 철회됐다. 2025년 2월 집행위원회가 철회를 발표했고 그해 8월 확정됐다. 이유는 이해관계자 간 합의 불가였다.^[3]

여기서 흔한 오해가 생긴다. "EU는 입증책임 전환을 이뤘고 한국은 못 했다"는 대조다. **이 전제는 틀렸다.** 그 지침은 지금 존재하지 않는 법이다.

그런데 다른 문이 열려 있었다

이야기가 여기서 끝나지 않는다.

같은 시기 EU는 **제조물책임지침**을 전면 개정했다. 이 지침은 소프트웨어를 "제품"의 정의에 넣었고, 결함 있는 AI 소프트웨어에 **무과실 책임**을 적용한다.^[4]

여기까지는 잘 알려진 이야기다. 그런데 그 안의 Article 10을 열어 보면 입증책임이 네 단계로 설계돼 있다.

1. 원칙 — 원고가 결함·손해·인과관계를 입증함.
2. **결함 추정** — 피고가 증거공개 의무를 어겼거나, 강행 안전요건을 위반했거나, 명백한 오작동이 있었으면 결함을 추정함.

3. 인과관계 추정 — 결함이 확정되고 손해가 그 결함과 전형적으로 부합하면 인과관계를 추정함.

4. 법원의 재량적 추정 — 기술적·과학적 복잡성 때문에 원고가 과도한 입증곤란을 겪고, 결함이나 인과관계의 개연성을 소명하면, 법원이 이를 추정해야 함.

네 번째 항의 요건 문언 — 기술적·과학적 복잡성으로 인한 과도한 입증곤란 — 은 AI 블랙박스성에 정확히 겹친다.

즉 채용 AI나 신용평가 AI가 "제품"에 해당하는 한, 철회된 지침이 하려던 일의 상당 부분을 이 조항이 이미 하고 있다.

다만 두 가지 한계가 남는다.

하나, 적용 대상이 제품에 한정된다. 소프트웨어는 제품에 포함됐지만, 순수하게 서비스로만 제공되고 어떤 소프트웨어 제품에도 해당하지 않는다고 다뤄질 여지가 있는 AI는 여전히 회색지대다. 법리가 아직 확립되지 않았다.

둘, 이것은 과실 추정이 아니라 결함·인과관계 추정이다. "AI Act 절차를 안 지켰다"는 사실만으로 과실을 추정하는 장치는 여전히 없다. 철회된 지침이 목표했던 것과는 법적 성격이 다르다.

한국은 어떤가. 인과관계 추정 조항이 없다. 입증책임 전환도 없다. 일반 불법행위 법리가 적용되고, 피해자가 스스로 인과관계를 입증해야 한다.

그런데 AI의 알고리즘과 학습 과정은 기업의 영업비밀이다. 법원을 통해 증거 개시를 청구할 수 있지만 기업이 영업비밀을 이유로 거부하면 소송이 거기서 막힌다.

보이지 않는 원인으로 받은 피해를 스스로 입증하라는 요구가 남는다.

세 번째 질문 — 구제

경로를 나란히 놓으면 이렇다.

	한국 A씨	EU A씨
이유 알기	AI 고지(사실만) 개인정보보호법 설명요구 (개인정보 처리 시간)	Art.86 의미 있는 설명 인간 검토 요청 가능
위반 확인	직접 조사 어려움	시장감시기관 조사 가능 피해자 민원 제기 가능
소송	인과관계 자기 입증 영업비밀로 증거 막힘 일반 불법행위 법리	제조물책임지침 Art.10(4) 법원의 재량적 추정 ※AI 책임지침은 철회됨
배상	손해액 스스로 증명	제품형 AI 무과실 책임 서비스형 AI는 회색지대

[필자 분석] 구조적 이점을 누가 깎는가

규제를 평가하는 시선은 둘이다.

기업의 시선 — 준수 비용이 얼마인가, 혁신을 막지 않는가.

피해자의 시선 — 피해를 입었을 때 실제로 구제받는가.

한국 인공지능기본법 논의에서 앞의 것은 많이 다뤄졌고 뒤의 것은 적게 다뤄졌다.

철회된 AI 책임지침의 출발점은 명확했다. AI의 불투명성은 만든 쪽의 구조적 이점이다. 기업은 자사 AI의 작동 방식을 알지만 피해자는 알 수

없다. 그 정보 비대칭이 책임 회피로 이어지지 않도록 입증책임을 재분배하는 것이 정의에 부합한다는 논리였다.

이 원칙은 지금도 유효하다. 틀린 것은 원칙이 아니라 그것을 법으로 만드는 과정이었다. 산업계의 반발과 합의 불가가 입법을 막았다.

그래서 정확한 대조는 이것이다.

EU는 제품형 AI에 무과실 책임을 확보했고, 그 안의 재량적 추정 조항이 기술적 복잡성을 이유로 한 입증 완화를 실질적으로 수행한다. 다만 순수 서비스형 AI에는 이 완화가 미치지 않을 수 있고, "의무 위반 곧 과실"이라는 장치는 EU에도 없다. **한국은 그 어느 것도 갖추지 못했다.**

EU도 완주하지 못했다는 사실이 중요하다. 이것은 한국이 뒤처졌다는 이야기가 아니라, **이 문제가 어렵다는 이야기다.** 어느 나라도 아직 깔끔한 답을 내놓지 못했고, 그래서 지금이 설계를 논의할 시점이다.

단기적으로는 고영향 AI 사업자 책무에 **설명 요구권**을 명시하는 일이, 중기적으로는 **인과관계 추정 조항**을 도입하는 일이 실질적 보완이 된다. 이것은 기업에 대한 부담이기만 한 것이 아니다. **구제 경로가 없는 기술은 결국 신뢰받지 못하고, 신뢰받지 못하는 기술은 팔리지 않는다.**

사고 다음에 남는 질문

피해자가 물을 곳이 없다면, 애초에 사고를 줄이는 쪽은 누가 지켜보는가.

다음 장에서 감시자를 본다.

실무자를 위한 체크박스

- ☐ 우리 AI가 개인에게 불이익한 결정을 내린다면, **결정 요소를 설명할 수 있는 로그가 남는가**. 설명 요구가 들어온 뒤 만들 수는 없다
- ☐ EU 기준으로 우리가 배포자라면 Art.86 설명 의무는 **우리에게** 온다. 모델 제공사에서 받는 정보로 그 설명이 가능한지 계약서에서 확인했는가
- ☐ 우리 AI 서비스가 제조물책임지침상 "제품"에 해당하는지 검토했는가. 해당하면 과실 책임과 법원의 재량적 추정이 함께 온다
- ☐ 개인정보를 처리하지 않는 자동화 결정이 있는가. 그렇다면 한국에서는 개인정보보호법의 설명요구권도 적용되지 않는다 — 규제 공백이지만 평판 리스크는 남는다
- ☐ 영업비밀을 이유로 증거 개시를 거부하는 것이 단기적으로는 유리해도, 재량적 추정 제도 아래서는 불리하게 작용할 수 있다는 점을 법무와 공유했는가

주(註)

1. AI Act Article 86. 고위험 AI가 관련한 결정에 적용되며, Annex III 일부 영역에 예외가 있다. 의무 주체는 배포자다.
2. 개인정보보호법 제37조의2. 완전히 자동화된 시스템으로 개인정보를 처리하여 이루어지는 결정이 권리·의무에 중대한 영향을 미치는 경우에 적용된다.
3. AI 책임지침 제안(COM(2022) 496). 2025년 2월 11일 집행위원회 2025 작업계획에서 철회가 발표되고 2025년 8월 확정됐다.
4. Directive (EU) 2024/2853. 2024년 12월 발효, 회원국 이행 기한 2026년 12월 9일. Article 4(1)이 소프트웨어를 제품에 포함시키고, Article 10이 입증책임 4단계를 규정한다. 특히 Article 10(4)는 기술적·과학적 복잡성으로 인한 과도한 입증곤란과 개연성 소명이 있으면 법원이 결함 또는 인과관계를 추정하도록 정한다.

11장. 누가 감시하는가 — AI Office와 과기정통부

이 장을 한 문장으로

EU는 40년 된 분산 감시 체계 위에 중앙 집행기구를 얹었고, 한국은 기초와 새 층이 서로 연결되지 않은 채 나란히 서 있다.

유럽연합에는 스물일곱 나라가 있고, 스물일곱 개의 시장감시기관이 있다.

독일 제품은 독일이 감시하고 프랑스 제품은 프랑스가 감시한다. 0부에서 본 체계다. 한 나라가 위험을 발견하면 경보망을 통해 나머지 전체로 퍼진다.

물리적 제품에 이 체계는 잘 작동했다. 독일 공장에서 나온 불량 토스터를 독일 당국이 발견하면 나머지 스물여섯 나라가 자국 시장에서 같은 제품을 확인한다.

그런데 EU는 AI에 대해 이 체계만으로는 부족하다고 판단했다. 그리고 전에 없던 기구를 만들었다.

분산 감시가 막히는 세 지점

첫째, 국경이 없다.

토스터는 독일 매장에서 팔린다. 물리적 실체가 있고 유통 경로를 추적할 수 있다.

그런데 AI 모델 하나가 클라우드에 올라가면 스물일곱 나라에서 동시에 작동한다. 서버는 아일랜드에 있고 개발사는 미국에 있고 사용자는 전 유럽에 흩어져 있다. 어느 나라 당국이 감시해야 하는가.

둘째, 범용 모델은 어디에도 속하지 않는다.

범용 AI 모델은 그 자체로 특정 제품이 아니다. 의료에 쓰이면 의료기기 영역이고 채용에 쓰이면 고위험 AI 영역이다. 그러나 모델 자체는 어느 분야 감시기관의 관할에도 깔끔히 들어가지 않는다.

40년 된 체계에서 감시기관은 **특정 법령의 특정 제품**을 맡도록 설계됐다. 모든 법령에 걸치는 범용 기술이라는 개념은 그 설계에 없었다.

셋째, 전문성의 벽.

토스터의 안전성을 보려면 전기공학 지식이면 된다. AI의 안전성을 보려면 머신러닝, 데이터 거버넌스, 편향 분석, 사이버보안, 그리고 적용 분야의 도메인 지식까지 필요하다.

스물일곱 나라가 모두 이 전문성을 갖추기는 현실적으로 어렵다. 작은 회원국의 감시기관이 최신 대형 모델을 독자적으로 평가할 수 있는가.

EU의 해법 — 기존 체계 위에 세 장치

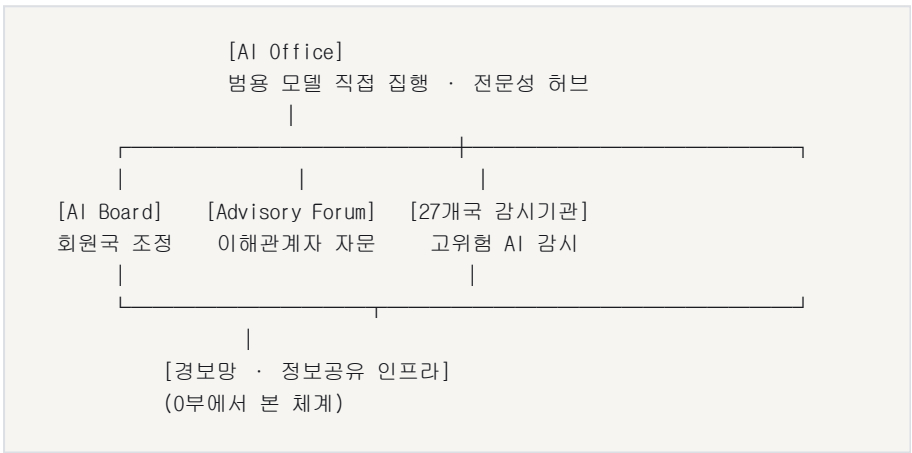
EU는 기존 감시 체계를 폐기하지 않았다. 그 위에 세 개를 얹었다.^[1]

AI Office는 집행위원회 안에 설치된 AI 전문 집행기관이다. 핵심은 하나다 — **범용 AI 모델에 대한 집행을 스물일곱 나라가 아니라 이곳이 직접 한다.** 평가하고, 정보를 요구하고, 위반 시 조치한다. 여기에 더해 유럽 차원의 AI 전문성 허브 역할과 표준화·샌드박스 조정을 맡는다.

AI Board는 회원국 대표들의 조정 기구다. 나라마다 다르게 집행되는 것을 막고 일관성을 보장한다.

Advisory Forum은 산업계·스타트업·시민사회·학계가 참여하는 자문기구다. 40년 된 체계에는 없던 구조인데, AI 규제가 기술과 사회와 윤리의 교차점에 있기 때문에 정부와 기업 밖의 목소리를 제도로 넣은 것이다.

구조를 그리면 이렇다.



핵심은 기존 체계가 그대로 돌아간다는 것이다. 스물일곱 나라는 여전히 자국 시장의 고위험 AI를 감시하고, 기존 시장감시 규정의 권한이 그대로 적용된다. 그 위에서 AI Office가 범용 모델만 직접 관할한다.

2장에서 본 "기존 건물에 새 층 올리기"가 거버넌스에서도 반복된다.

그 안의 두 가지 논쟁

이 설계가 순탄하게 나온 것은 아니다.

보충성 원칙과의 긴장. EU는 전통적으로 "회원국이 할 수 있는 일은 회원국에 맡긴다"는 원칙을 따른다. 시장감시가 대표적이다. 그런데 AI Office는 범용 모델 집행을 브뤼셀에 집중시켰다. EU 제품안전 역사에서 전례가 드문 중앙 집행이다.^[2]

독립성 문제. AI Office는 별도 법인격을 가진 독립 기관이 아니라 집행위원회 내부 조직이다. 의사결정이 빠르다는 장점이 있고, 정치적 판단에 영향받을 수 있다는 우려가 있다. AI 산업 육성과 AI 규제 집행을 같은 조직 안에서 하는 것이 적절한가 — 이 질문은 아직 열려 있다.

한국은 어떻게 되어 있는가

한국은 과학기술정보통신부 한 곳이 총괄한다. EU의 분산 모델과 정반대 지점에서 출발한다.

법 제12조가 설립 근거를 둔 인공지능안전연구소를 AI Office의 한국판으로 보는 시각이 있다. 그러나 권한의 성격이 다르다.

인공지능안전연구소의 임무는 위험 정의·분석, 평가 기준과 방법 연구, 프론티어 모델 이행결과 접수·평가 지원 등 연구·평가·자문 중심이다. AI Office가 가진 정보 요청권, 시정조치권, 과징금 부과권 같은 직접 집행권은 법령에 명시돼 있지 않다.^[3]

집행은 과학기술정보통신부장관에게 남아 있다. 제40조의 사실조사와 시정명령이 그 수단이다.

여기서 자주 혼동되는 조직이 하나 있다. 국가인공지능전략위원회(대통령 소속)는 이 장이 다루는 감시·집행 기구가 아니다. 기본계획 심의, 부처 간 조율, 고영향 AI 규율 방향 같은 정책 결정을 맡는다. 위원 정수는 최초 발의 시 45명이었다가 2026년 1월 개정으로 60명 이내로 늘었다.^[4] 고영향

AI 확인을 자문하는 **전문위원회**(50명 이내)는 이 전략위원회와 별개로 과학기술정보통신부 산하에서 운영된다 — 이름이 비슷해 둘을 같은 조직으로 오해하기 쉽다.

[필자 분석] 단순함인가 빈칸인가

AI Office의 탄생은 기존 체계의 진화인 동시에 한계의 고백이다.

40년 동안 그 체계는 "분산 감시, 상호 인정"으로 작동했다. 회원국이 각자 감시하되 하나의 기준으로 통일한다. 이 모델이 성공한 이유는 **제품이 물리적이고 시장이 지역적이었기 때문이다**.

AI는 두 전제를 모두 깼다. 제품은 소프트웨어이고 시장은 즉시 전 지구적이다. 그래서 EU는 분산 모델로 감당할 수 없는 영역에 한해 중앙 집행을 택했다. 고위험 AI는 여전히 각국이 보되, **범용 모델만큼은 브뤼셀이 직접 본다**.

한국과 비교하면 흥미로운 대칭이 나온다. 한국은 처음부터 한 부처가 총괄하니 EU가 애써 해결한 조정 문제 자체가 없다. 이것만 보면 한국이 단순하고 효율적이다.

그런데 더 근본적인 차이가 따로 있다.

EU는 AI Office를 기존 시장감시 체계 위에 올렸다. 기초가 있는 건물에 층을 더한 것이다. 아래층에는 적합성평가를 수행할 기관이 있고, 그 기관을 검증할 인정 체계가 있고, 문제를 잡아낼 감시망이 있다.

한국은 기초와 새 층이 서로 연결되지 않은 채 나란히 서 있다. 제품안전 시장감시 체계와 인공지능기본법 사이에 둘을 잇는 연결 규범이 없다. 2장에서 확인한 "기초 없이 새 층만 올린 형국"이 거버넌스에서 다시 확인되는 셈이다.

그래서 이렇게 정리할 수 있다. EU의 복잡함은 40년간 쌓아 온 체계 위에서의 복잡함이다. 한국의 단순함은 효율일 수도 있고, 아직 아무것도 짓지 않아서 생긴 빈칸일 수도 있다.

둘을 가르는 것은 조직도가 아니라 그 아래 무엇이 깔려 있는가다.

3부를 마치며

3부에서 우리는 사고 이후를 봤다.

잘못됐을 때 피해자가 물을 곳이 있는가(10장), 그리고 평소에 누가 지켜보는가(11장).

두 장의 답이 같은 방향을 가리킨다. 한국법에는 조문이 있지만 그 조문을 실제로 작동시킬 아래층이 아직 얇다. 설명요구권 조항이 없고, 인과관계 추정이 없고, 집행권을 가진 전문기관이 없다.

그리고 EU도 완주하지는 못했다. 서비스형 AI의 책임 공백은 그쪽에도 남아 있고, 중앙 집행기구의 독립성 문제도 아직 답이 없다.

한줄 코멘트. 감시 체계를 비교할 때 봐야 할 것은 조직도가 아니라 그 조직이 딛고 선 바닥이다 — EU의 복잡함은 40년 위에 쌓인 복잡함이고, 한국의 단순함은 효율일 수도 빈칸일 수도 있다.

이제 시선을 바꿀 차례다. 지금까지는 두 법이 어떻게 생겼는지를 봤다. 4부에서는 그 법 앞에 선 기업이 무엇을 해야 하는지를 본다.

실무자를 위한 체크박스

- 우리가 범용 모델 제공자라면 EU에서 상대할 곳은 회원국 당국이 아니라 **AI Office**다. 창구가 다르면 대응 방식도 다르다
- 고위험 AI를 공급한다면 상대는 **각 회원국 감시기관**이다. 주요 진출국의 담당 기관을 파악해 두었는가
- 한국에서 집행 주체는 과학기술정보통신부장관이고, 인공지능안전연구소는 평가·자문 역할이다. **문서가 어디서 오는지에 따라 대응 성격이 다르다**
- 제40조 사실조사에 대응할 사내 절차가 있는가. 현장조사와 자료 제출 요구가 가능하냐
- Advisory Forum처럼 이해관계자가 제도적으로 참여하는 창구를 **의견 개진 기회**로 활용할 계획이 있는가 (15장 참조)

주(註)

1. AI Act Article 64(AI Office), Article 65-66(AI Board), Article 67(Advisory Forum). 범용 AI 모델에 대한 AI Office의 직접 집행 권한은 Article 88-94에 있다.
2. 의약품이나 화학물질 분야에는 유사한 중앙 기관 선례가 있으나, 제품안전 체계 안에서는 AI Office가 최초로 가까운 중앙 집행 사례다.
3. 인공지능기본법 제12조(인공지능안전연구소) 및 시행령. 임무는 위험 정의·분석, 평가 기준·방법 연구, 법 제32조 이행결과 관련 지원 등 연구·평가·자문 중심이며, 정보 요청·시정조치·과징금 부과 같은 직접 집행권은 명시돼 있지 않다.
4. 인공지능기본법 제7조(국가인공지능전략위원회). 위원 정수는 최초 발의 시 45명이었다가 2026년 1월 개정으로 60명 이내로 확대됐다. NIA 채은선, "인공지능기본법의 주요 내용" 세미나(AI 제품안전 혁신 TF 제6차 전체 회의, 2026-08-06) 인용. 상세는 위키 국가인공지능전략위원회 참조.

PART 4

4부 — 기업의 눈: 규제를 사업으로 번역하기

이 부는 실무자를 위한 것이다. 순서는 그 자리에서 실제로
일이 흘러가는 순서를 따랐다.

12장. 우리 서비스는 어느 칸에 들어가는가

이 장을 한 문장으로

두 법 모두 "당신의 서비스가 무엇인가"를 먼저 묻지만, EU는 나무를 타고 내려가게 하고 한국은 목록에서 찾게 한다.

해마다 봄이면 같은 장면이 반복된다. 서류를 앞에 두고, 지난해 들어온 돈을 어느 칸에 적을지 고민한다. 근로소득인가, 사업소득인가, 기타소득인가.

같은 금액인데도 어느 칸에 들어가느냐에 따라 세율이 달라진다. 내야 할 서류가 달라지고, 빠뜨렸을 때 치를 대가도 달라진다.

그리고 그 판단은 1차적으로 본인이 한다. 세무서가 먼저 찾아와 "당신 소득은 이것입니다"라고 정해 주지 않는다. 스스로 칸을 고르고, 그 선택이 맞았는지는 나중에 검증받는다.

AI 규제 첫 단계가 정확히 이 구조다. 어느 칸에 들어가는지가 정해지는 순간, 해야 할 일의 목록과 마감일과 비용이 함께 정해진다. 그리고 그 칸을 고르는 것은 우선 사업자 자신이다.

EU는 나무를 타고 내려간다

EU AI Act에서 분류는 위에서부터 순서대로 걸러 내려가는 구조다. 위쪽에서 걸리면 아래는 볼 필요가 없다.

1. **금지 목록에 걸리는가.** 걸리면 끝이다. EU 시장에 낼 수 없다. 열 가지 항목은 4장에서 다뤘음.
2. **Annex I 제품의 안전부품인가.** 의료기기·기계·완구처럼 이미 EU 제품안전법이 규율하는 물건에 AI가 들어갔고, 그 AI가 안전 기능을 맡는 경우임. 걸리면 그 제품법과 AI Act가 함께 적용됨.
3. **Annex III 여덟 영역에 해당하는가.** 제품의 부품이 아니라 독립된 소프트웨어로서, 채용·대출·교육·생체인식·공공서비스·법집행·이민·사법 중 하나에 쓰이는 경우임.
4. **투명성 의무 대상인가.** 챗봇·합성 콘텐츠 생성·딥페이크·감정인식. 고위험이 아니어도 알릴 의무가 붙음(6장).
5. **범용 AI 모델인가.** 누적 연산량 10^{25} FLOPs를 넘으면 시스템적 위험 GPAI로 추가 의무가 붙음(8장).
6. **어디에도 해당하지 않는가.** 그러면 AI Act상 의무는 없다.

이 순서에서 가장 중요한 갈림길은 두 번째와 세 번째 사이다. 몸을 가진 AI냐, 몸이 없는 AI냐. 같은 알고리즘이라도 의료기기에 들어가면 의료기기규정과 AI Act를 동시에 통과해야 하고, 웹 서비스로 제공되면 AI Act만 상대하면 된다.

2026년에 바뀐 갈림길

그런데 이 두 번째 관문이 최근 좁아졌다. 2026년 7월 발효한 개정에서 "안전부품"의 정의가 줄었다.^[1]

이전에는 조화법령이 적용되는 제품 안에 AI가 들어가 있으면 안전부품으로 볼 여지가 넓었다. 개정 후에는 그 AI의 의도된 목적이

사람이나 재산의 위험을 예방하거나 줄이는 것일 때만 안전부품이다.

무엇이 빠졌는지가 실무적으로 중요하다. 사용자 지원, 성능 최적화, 서비스 효율, 자동화, 편의, 품질 관리. 이런 목적만으로 쓰이는 AI는 조화법령 제품 안에 들어 있어도 안전부품이 아니다. 냉장고의 식재료 추천 기능, 세탁기의 세탁 코스 자동 선택 같은 것들이 여기 해당한다.

다만 예외가 붙는다. **고장 나거나 오작동했을 때 건강과 안전을 위협할 수 있으면 여전히 안전부품이다.** 목적이 편의였더라도 결과가 위험하면 걸린다.

한 가지 덧붙일 것이 있다. 제3자 적합성평가가 요구되는 이유가 건강·안전이 아닌 다른 사유뿐이라면 — 예컨대 전파 간섭 같은 것 — 고위험 조건을 채우지 못한다. 인증을 받는다는 사실 자체가 고위험을 뜻하지는 않는다는 뜻이다.

한국은 목록에서 찾는다

한국 인공지능기본법의 판정은 분기가 아니라 대조다. 법이 정한 **열 개 영역**을 펼쳐 놓고, 우리 서비스가 그 안에 있는지 확인한다.^[2]

에너지 공급, 먹는물, 보건의료, 의료기기, 원자력, 범죄수사 생체인식, 채용·대출심사, 교통, 국가기관 의사결정, 유아·초중등 교육.

여기 해당하면 고영향 인공지능이고, 제34조의 여섯 가지 책무가 따라붙는다. 해당하지 않으면 그 책무는 없다.

EU와 달리 "**제품에 붙은 AI**"와 "**독립 소프트웨어 AI**"를 나누는 갈림길이 없다. 의료기기에 들어간 AI든 독립 진단 소프트웨어든 같은 목록의 같은 항목으로 본다. 제품안전법 계보와 연결되지 않은 구조가 여기서 드러난다.

판정 절차도 다르다. 사업자가 스스로 검토하는 것은 의무다. 그러나 과학기술정보통신부장관에게 확인을 요청하는 것은 재량이다.^[3] 스스로 판단하고 그대로 갈 수도 있고, 확인을 받아 둘 수도 있다.

그 밖에 두 가지 관문이 더 있다. 생성형 AI라면 표시 의무가 붙고(제31조), 누적 연산량이 10^{26} FLOPs를 넘으면 안전성 확보 의무가 붙는다(제32조).

우리 서비스는 하나가 아니다

실무에서 이 판정이 어려운 진짜 이유는 조문이 복잡해서가 아니다.

하나의 서비스 안에 여러 개의 AI가 들어 있기 때문이다.

고객 상담 플랫폼 하나를 생각해 보자. 대화를 처리하는 언어 모델이 있고, 상담원 배정을 최적화하는 모델이 있고, 통화 내용에서 불만 강도를 읽는 모델이 있고, 신규 상담원을 뽑을 때 지원서를 걸러 주는 모델이 있다.

네 개의 모델은 서로 다른 칸에 들어간다. 대화 처리는 투명성 의무 대상이고, 배정 최적화는 아무 데도 걸리지 않으며, 감정 추론은 사용 맥락에 따라 금지이거나 고위험이고, 채용 선별은 EU에서는 Annex III 고위험이고 한국에서는 고영향이다.

"우리 회사는 고위험 AI 기업인가"는 답할 수 없는 질문이다. 답할 수 있는 것은 "이 기능은 어느 칸인가"뿐이다.

[필자 분석] 같은 자기 판정, 다른 사후 검증

두 법 모두 1차 판정을 사업자에게 맡긴다. 이것은 0-2에서 본 40년짜리 전통의 연장이다. 유럽 제품안전 체계는 애초에 국가가 모든 제품을 미리

검사하는 방식을 포기하고, 제조사가 스스로 선언하되 그 선언에 책임을 지우는 쪽을 택했다.

그런데 자기 판정을 맡긴 뒤 무엇을 요구하는가에서 두 법이 갈린다.

EU는 판정 자체를 문서로 남기게 한다. 고위험 목록에 해당하지만 실질적 위험이 없다고 판단해 예외를 주장하려면, 그 판단 근거를 시장에 내놓기 전에 작성해 두어야 하고 당국이 요구하면 제출해야 한다.^[4] 판정이 틀렸는지를 나중에 따져 볼 수 있는 기록이 남는다.

한국은 검토를 의무로 두되 그 결과를 남기라는 요구가 조문에 없다. 확인 요청은 재량이므로, 사업자가 "우리는 고영향이 아니다"라고 판단하고 넘어가면 그 판단의 흔적이 어디에도 남지 않는다.

같은 자기 판정이지만 되짚을 수 있는 정도가 다르다. 그리고 이 차이는 사고가 났을 때 드러난다. 문서가 있으면 판단의 잘잘못을 다룰 수 있고, 없으면 다룰 재료 자체가 없다.

칸을 고르는 일이 끝나면 다음 질문이 온다. 그래서 무엇을 만들어야 하는가.

실무자를 위한 체크박스

- ☐ 우리 조직이 만들거나 사용하는 AI 기능을 **기능 단위로** 나열했는가 — 제품 단위가 아니라 기능 단위여야 한다
- ☐ 각 기능마다 EU 여섯 관문과 한국 네 관문을 따로 통과시켜 봤는가
- ☐ 조화법령 제품에 들어간 AI라면, 그 AI의 **의도된 목적**이 안전인지 편의인지 문서로 정리했는가 (2026년 개정으로 갈리는 지점)
- ☐ 고위험·고영향이 아니라고 판단한 기능이 있다면, 그 판단의 근거를 **지금** 문서로 남겨 두었는가
- ☐ 판정 결과를 제품 출시·업데이트 절차에 연결해 두었는가 — 기능이 추가되면 판정도 다시 해야 한다(5장)

주(註)

1. Regulation (EU) 2026/1744(Digital Omnibus on AI) Article 1(4)(a)·1(8), AI Act Article 3(14) 개정 및 Article 6(1a)~(1c) 신설. 발효 2026-07-27. 안전 기능은 공급자가 정하는 의도된 목적이며, 조화법령 제품에 통합돼 있다는 사실만으로는 안전 기능을 수행한다고 보지 않는다(전문 (7)).
2. 인공지능 발전과 신뢰 기반 조성 등에 관한 기본법 제2조제4호 가목~차목. 카목은 대통령령 위임이며 현재 시행령에 추가 지정은 없다.
3. 인공지능기본법 제33조(고영향 인공지능의 확인). 사업자의 사전 검토는 의무이나, 장관에 대한 확인 요청은 재량이다.
4. AI Act Article 6(3) — Annex III 해당 AI라도 중대한 위험을 야기하지 않는 경우의 예외. 공급자는 시장 출시 전에 그 평가를 문서화해야 하며, 국가 관할당국은 그 평가의 제출을 요구할 수 있다. Regulation (EU) 2026/1744 전문 (22)이 이 문서화 의무가 유지됨을 확인하면서 EU 데이터베이스 등록 항목만 간소화했다.

13장. 두 법을 동시에 맞추는 최소 공통 작업

이 장을 한 문장으로

두 법을 각각 맞추면 두 번 일하게 된다. 겹치는 것을 먼저 만들면 한 번의 작업이 두 곳을 덮는다.

여행 가방을 싸 본 사람은 안다. 유럽과 일본을 연달아 가야 할 때, 챙기는 것은 대부분 같다. 옷도 같고 세면도구도 같고 노트북도 같다.

다른 것은 콘센트뿐이다.

그래서 노트북을 두 대 사지 않는다. 어댑터를 두 개 챙긴다. **본체는 하나로 만들고, 나라마다 다른 부분만 따로 붙이는 것이 짐을 줄이는 유일한 방법이다.**

두 AI법 대응도 같은 원리로 접근할 수 있다. 그런데 실제 현장에서는 반대로 하는 경우가 많다. EU 대응 프로젝트를 따로 열고, 국내 대응 프로젝트를 따로 연다. 두 팀이 각자 문서를 만들고, 그 문서가 서로를 모른다.

무엇이 겹치고 무엇이 다른가

두 법의 요구사항을 나란히 놓으면 뚜렷한 규칙이 하나 보인다.

거의 모든 항목에서 EU 쪽 요구가 더 높거나 같다. 낮은 경우는 없다시피 하다.

이것이 실무적으로 뜻하는 바는 분명하다. EU 기준으로 한 번 만들면 한국 요건은 대체로 함께 충족된다. 반대 방향은 성립하지 않는다. 한국 기준으로 만든 문서를 들고 EU에 가면 거의 전부 다시 만들어야 한다.

몇 가지만 짚어 보자.

- **리스크 관리** — EU는 생애주기 전체에 걸친 반복 순환 시스템을 요구하고(Art.9), 한국은 위험관리방안 수립을 요구함(제34조1호). 앞의 것을 만들면 뒤의 것은 그 안에 들어 있음.
- **인적 감독** — EU Art.14와 한국 제34조4호가 요구하는 바가 거의 같음. 여기는 애초에 겹침.
- **문서 보관** — EU는 10년, 한국은 5년임. 긴 쪽에 맞추면 됨.
- **표시·고지** — EU는 기계가 읽을 수 있는 형식을 의무로 요구하고, 한국은 선택으로 둠(6장). 기계판독으로 만들면 양쪽이 덮임.

다만 어댑터가 필요한 자리도 있다. 관할이 다르면 본체를 공유할 수 없는 항목이다.

국내대리인·수권대리인 지정은 각 관할에 따로 해야 한다. 표시 문구의 언어도 따로다. 신고와 등록의 상대방도 다르다. 이런 것들은 아무리 잘 설계해도 두 번 해야 한다.

순서가 비용을 정한다

무엇을 먼저 하느냐가 전체 비용을 크게 좌우한다. 아래 순서를 권한다.

1. **AI 인벤토리부터 만든다.** 우리 조직이 만들고, 사고, 쓰는 AI가 각각 몇 개이고 어디에 붙어 있는지의 목록임. 이것이 없으면 다음 단계 전부가 공중에 뜬.

2. 기능 단위로 분류한다. 12장의 판정을 인벤토리의 각 줄에 적용함. 이 단계에서 대부분의 항목이 "아무 데도 걸리지 않음"으로 정리되고, 남은 것은 생각보다 적음.
3. 남은 것에만 리스크 관리 문서를 만든다. EU Art.9 기준으로 설계함. 한국 제34조1호는 자동으로 덮임.
4. 기술문서와 로그를 같은 틀로 만든다. EU Annex IV의 열 개 항목이 사실상 업계 표준 목차 역할을 함. 한국에는 대응 조문이 없지만, 만들어 두면 제34조5호 문서 보관 의무가 함께 충족되고 사고가 났을 때 방어 재료가 됨(10장).
5. 표시 체계를 기계판독 기준으로 한 번에 만든다. 2026년 12월 2일이 EU 쪽 경과조치 종료일임(부록 B).
6. 사고 기록 절차를 만든다. EU는 중대사고 보고를 의무로 요구하고 한국에는 대응 조문이 없음. 그러나 제조물책임 대응 관점에서는 한국에서도 필요함.

이 순서의 핵심은 1번과 2번이다. 인벤토리와 분류를 건너뛰고 3번부터 시작하면, 걸리지도 않는 기능에 문서를 만들게 된다. 실무에서 가장 흔한 낭비가 여기서 생긴다.

먼저 오는 것은 한국이다

여기서 역설이 하나 나온다.

지금까지 "EU 기준으로 만들라"고 했는데, 정작 먼저 지켜야 하는 것은 한국 쪽이다.

한국 인공지능기본법은 2026년 7월 21일에 전면 시행됐다. 고영향 AI 사업자의 여섯 가지 책무는 이미 작동하고 있다. 반면 EU의 Annex III 고위험 AI 요건은 2027년 12월 2일로, 제품 내장형은 2028년 8월 2일로 미뤄졌다(9장·부록 B).

한국이 EU보다 1년 4개월 앞서 있다.

그래서 앞의 조언이 모순처럼 들린다. 당장 필요한 것은 한국 요건인데 설계는 EU 기준으로 하라니, 급하지 않은 쪽에 먼저 힘을 쓰라는 말처럼 읽힌다.

모순이 아니다. 두 가지를 나누면 된다.

- **지금 제출해야 하는 것** — 한국 기준을 채우면 됨. 마감은 이미 지났음.
- **지금 설계해야 하는 것** — EU 기준이어야 함. 마감은 1년 넘게 남았지만, 나중에 다시 만들면 지금 한 작업이 버려짐.

실무에서는 이렇게 된다. 리스크 관리 체계는 EU 제9조 수준으로 설계하되, 당장 제출하는 문서는 그중 한국 제34조제1호가 요구하는 부분만 뽑아서 낸다. 틀은 크게 짜고, 지금 채울 칸만 먼저 채운다.

거꾸로 하면 대가가 크다. 한국 기준에 맞춰 작게 만들어 두고 나중에 넓히려 하면, 넓히는 일이 처음부터 만드는 일과 비용이 거의 같아진다. 문서의 목차 구조 자체가 다르기 때문이다.

문서가 아니라 조직이 비용이다

여기까지는 만들어야 할 산출물의 이야기였다. 그런데 진짜 비용은 다른 곳에 있다.

만든 것을 계속 살아 있게 유지하는 일.

EU AI Act 제17조가 요구하는 품질관리시스템은 문서 한 벌이 아니다. 그 문서를 누가 갱신하고, 모델이 바뀌면 누가 다시 검토하고, 그 검토가 실제로 이뤄졌음을 무엇으로 증명하는가에 관한 체계다. 5장에서 본 "검진과 주치의"의 차이가 여기서 다시 나온다.

한국법에는 이에 해당하는 요건이 없다. 그래서 국내 사업만 하는 동안에는 이 비용이 눈에 띄지 않는다.

EU에 나가는 순간 드러난다. 그리고 그때 발견하면 늦다. 조직과 절차를 만드는 데는 문서를 만드는 것보다 훨씬 긴 시간이 걸리기 때문이다.

[필자 분석] "최소 공통"이라는 말의 함정

이 장의 제목에 쓴 "최소 공통"이라는 표현에는 오해의 소지가 있다. 겹치는 것만 하면 두 법을 다 맞출 수 있다는 뜻으로 읽힐 수 있기 때문이다.

그렇지 않다. 겹치는 것만 하는 것이 아니라, 겹치는 것부터 하는 것이다.

두 법의 요구가 겹치는 영역은 대체로 두 입법자가 모두 "이건 반드시 필요하다"고 판단한 부분이다. 리스크를 관리할 것, 사람이 개입할 수 있게 할 것, 기록을 남길 것. 서로 다른 철학에서 출발한 두 법이 같은 결론에 도달했다면, 그 항목은 규제 대응 이전에 **제품을 제대로 만드는 일**에 가깝다.

반대로 한쪽에만 있는 항목은 그 법의 고유한 관점을 담고 있다. EU에만 있는 데이터 거버넌스와 자동 로그는 "AI의 위험은 데이터와 추적 불가능성에서 온다"는 판단의 산물이고, 한국에만 있는 우선구매 확인은 "시장을 만들어 유도한다"는 판단의 산물이다.

겹치는 부분은 제품의 문제이고, 갈리는 부분은 시장의 문제다. 앞의 것을 먼저 하는 편이 순서상 자연스럽고, 실제로 비용도 덜 든다.

체계를 만들어 두면 그것으로 끝나는가. 그렇지 않다. 규제는 계속 움직이고, 움직일 때마다 만들어 둔 것을 고쳐야 한다. 그러면 무엇이 언제 움직이는지는 어떻게 아는가.

실무자를 위한 체크박스

- ☐ SI 인벤토리가 문서로 존재하는가 — 없다면 이것이 첫 번째 과제다
- ☐ 문서를 만들 때 **높은 쪽(대체로 EU) 기준**으로 설계하고 있는가, 아니면 각 법마다 따로 만들고 있는가
- ☐ 관할별로 반드시 나뉘야 하는 항목(대리인 지정·표시 언어·신고 상대방)을 별도로 식별해 두었는가
- ☐ 만든 문서를 **누가 언제 갱신하는지**가 정해져 있는가 — 이것이 EU Art.17이 실제로 묻는 것이다
- ☐ 모델 변경·기능 추가 시 분류와 문서를 다시 검토하는 절차가 개발 프로세스에 들어가 있는가

14장. 정책 환경을 읽는 법 — 어디를 보고 무엇을 기다리나

이 장을 한 문장으로

규제 대응의 절반은 이미 나온 법을 읽는 일이고, 나머지 절반은 아직 나오지 않은 것을 기다리는 일이다.

내일 비가 올지 알고 싶으면 기상청을 본다. 기상청은 한 곳이고, 예보는 정해진 시간에 나오고, 형식도 일정하다. 우산을 챙길지 말지는 그 하나를 보고 정할 수 있다.

규제에는 기상청이 없다.

관보가 있고 부처 발표가 있고 표준화 기구의 진행 상황이 있고 업계 보도가 있다. 이 창구들은 서로 다른 주기로, 서로 다른 형식으로, 서로 다른 확실성을 가지고 정보를 내놓는다. 그것을 하나의 예보로 엮는 일은 스스로 해야 한다.

어디를 볼 것인가

지켜볼 곳은 양쪽에 네 군데씩이다.

EU 쪽 네 곳.

1. 관보(Official Journal of the EU) — 법이 확정되는 유일한 장소임. 여기 실리기 전에는 어떤 것도 최종이 아님.

2. **AI Office 발표** — 가이드라인, 행동강령(Code of Practice), 템플릿. 법은 아니지만 무엇을 어떻게 해야 하는지의 실무 기준이 여기서 나옴.
3. **조화표준 진행 상황** — CEN/CENELEC JTC 21이 AI Act용 유럽 표준을 만들고 있음. 뒤에서 따로 다룸.
4. **Safety Gate** — 실제로 무엇이 문제가 되어 시장에서 회수됐는지가 쌓이는 곳임. 조문이 아니라 사고를 보는 창구.

한국 쪽 네 곳.

1. **관보** — 법률·시행령·고시가 확정되는 곳임.
2. **과학기술정보통신부 고시 행정예고** — 이 목록에서 실무적으로 가장 중요한 창구일 수 있음. 고시가 확정되기 전 예고 단계이고, **이때 의견을 낼 수 있기 때문임**.
3. **국가인공지능전략위원회 의결 사항** — 정책 방향이 먼저 드러나는 자리.
4. **인공지능안전연구소** — 위험 평가 기준·방법론이 나오는 곳.

같은 뉴스가 아니다

여기서 이 장의 핵심으로 들어간다. 위 창구들에서 나오는 정보는 **확실성의 등급이 다르다**. 이것을 섞으면 대응이 흔들린다.

네 등급으로 나눠 보자.

- **확정** — 관보에 실렸음. 번호가 붙었고 시행일이 정해졌음.
- **임박** — 절차상 남은 것이 형식뿐임. 행정예고를 마친 고시, 최종 표결만 남은 법안.
- **방향** — 초안이 공개됐거나 의견수렴 중이거나 정치적 합의에 이르렀음. 내용은 바뀔 수 있음.

- **관측** — 언론 보도, 업계 전망, 전문가 해설. 근거를 따라 올라가면 위 셋 중 하나에 닿아야 함.

2026년의 EU 개정이 이 구분을 잘 보여 준다. 그해 5월에 이사회와 의회가 정치적 합의에 이르렀고, 6월에 의회가 1차독회 채택을 했고, 7월 8일에 최종 채택이 이뤄졌으며, 7월 24일에 관보에 실려 **Regulation (EU) 2026/1744**라는 번호를 받았다.

이 녀 달 동안 시중에는 "고위험 AI 적용이 2027년 12월로 연기됐다"는 문장이 돌아다녔다. 그 문장의 내용은 결과적으로 맞았다. 확정본과 대조해보면 낱자도 수치도 합의 단계 요약과 일치했다.

그러나 4월에 그 문장을 근거로 사업 계획을 바꾼 조직과, 관보를 기다린 조직은 다른 위험을 안고 있었다. 내용이 맞았던 것은 결과이지 보장이 아니었다.

부정형 사실이 가장 빨리 낡는다

모니터링에서 놓치기 쉬운 함정이 하나 있다.

"아직 없다"는 사실은 "있다"는 사실보다 빨리 낡는데, 낡아도 티가 나지 않는다.

새 법이 생기면 뉴스가 된다. 그러나 어제까지 없던 고시가 오늘 생긴 것은, 그것을 기다리던 사람에게만 뉴스다. 기다리지 않던 사람의 문서에는 여전히 "미제정"이라고 적혀 있고, 그 문장은 아무 경고 없이 틀린 문장이 된다.

대응은 단순하다. "아직 없다"고 적을 때는 언제 확인했는지를 함께 적는다. 그러면 그 문장을 읽는 사람이 스스로 유효기간을 판단할 수 있다.

표준을 기다린다는 것

EU 쪽에서 가장 기다릴 만한 것은 조화표준이다. 왜 그런지는 설명이 필요하다.

AI Act 제40조는 조화표준을 지키면 요건을 충족한 것으로 추정해 준다. 이 추정이 있으면 "무엇을 어디까지 하면 되는가"가 확정된다. 없으면 각자 알아서 판단해야 하고, 그 판단이 맞는지는 나중에 당국이 정한다.

표준을 만드는 곳은 CEN/CENELEC JTC 21이다. 2021년 6월에 설립됐고 300명 넘는 전문가가 참여한다. 집행위원회가 표준화 위임명령을 내렸고, 그에 따라 개별 표준이 개발되고 있다.

품질관리시스템에 관한 prEN 18286이 2025년 10월 공개 심의에 부쳐졌다. 제17조 품질관리시스템 요건에 대응하는, AI Act용으로는 첫 표준이다.^[1]

이름 자체가 등급을 알려 준다. 앞의 pr은 아직 제안 단계라는 표시다. 표결과 채택을 거쳐 정식 유럽표준이 되면 이 두 글자가 떨어지고 EN 18286이 된다. 문서 번호만 보고도 확정인지 방향인지 구분할 수 있는 셈이다.

집필 시점인 2026년 8월 초까지 그 두 글자는 아직 붙어 있다. 관보 게재도 확인되지 않는다. 게재되는 순간 준수 추정이 작동하고, 그 전까지 각 조직이 자체적으로 세워 둔 기준은 재검토 대상이 된다. 이것이 부록 B의 대기 목록에 이 항목을 넣어 둔 이유다.

한국 쪽에는 준수 추정 메커니즘 자체가 없다. 대신 고시·가이드라인이 대기 목록이다. 우선구매 확인 절차는 2026년 7월에 나왔고, 연산량 산정방식·투명성 표시 예외·영향평가 세부 기준은 아직이다(부록 B). 관련 가이드라인의 작성 주체는 나뉜다 — 고영향 AI 판단 가이드라인은

한국지능정보사회진흥원(NIA)이, 고영향 AI 사업자 책무 고시·가이드라인은 한국정보통신기술협회(TTA)가 만든다.^[2] 가이드라인이 나왔다고 그 근거인 의무 고시까지 제정된 것은 아니다 — 둘을 구분해서 추적해야 한다.

[필자 분석] 모니터링의 목적은 수집이 아니다

정보를 모으는 것 자체는 목적이 아니다. 모니터링의 목적은 결정을 내릴 시점을 확보하는 것이다.

조화표준이 언제 나오는지를 왜 알아야 하는가. 그것이 나오면 대응 방식이 확정되기 때문이다. 그 전까지 조직은 둘 중 하나를 하게 된다. 최대한으로 대비하거나(과잉), 나올 때까지 미루거나(과소). 어느 쪽이든 비용이다.

나올 시점을 알면 세 번째 선택지가 생긴다. **표준안의 방향에 맞춰 지금 할 수 있는 것만 하고, 확정되면 나머지를 채우는 것.** 이 선택지는 시점을 아는 조직에만 열린다.

한 걸음 더 나아갈 수도 있다. 조화표준은 완성된 채로 하늘에서 내려오지 않는다. 사람들이 회의실에 모여 만든다. **그 회의실에 앉을 수 있다면, 기다리는 대신 쓰는 쪽이 된다.**

지켜보는 것과 개입하는 것은 다르다. 다음 장은 그 차이에 관한 이야기다.

실무자를 위한 체크박스

- ☐ EU 네 곳·한국 네 곳의 확인 주기를 정해 두었는가 — 담당자와 주기가 없으면 모니터링이 아니라 우연이다
- ☐ 사내 규제 자료에서 **확정·임박·방향·관측**을 표기로 구분하고 있는가
- ☐ "아직 없다"고 적힌 문장마다 **확인 일자**가 붙어 있는가
- ☐ 과기정통부 고시 행정예고를 확인하고 있는가 — 확정 후에 아는 것과 예고 단계에서 아는 것은 대응 여지가 다르다
- ☐ prEN 18286을 비롯한 조화표준 게재 여부를 대기 목록에 올려 두었는가(부록 B)

주(註)

1. CEN/CENELEC JTC 21은 2021년 6월 설립된 유럽 AI 표준 개발 전담 기술위원회이며, 집행위원회 표준화 위임명령 C(2023)3215에 따라 AI Act 대응 표준을 개발한다. prEN 18286(AI Quality Management System for EU AI Act Regulatory Purposes)은 AI Act Article 17 품질관리시스템에 대응하는 첫 AI Act 표준으로, **공개 심의(public enquiry)가 2025-10-30부터 2025-12-27까지 진행됐다**. 회원국 표결·채택과 관보 게재를 거쳐야 정식 유럽표준이 되며 Article 40의 준수 추정이 작동한다. **2026-08-05 조화 기준 여전히 draft(prEN) 단계이며 관보 게재는 확인되지 않았다**. 상세는 eu-ai-act-iso-iec-표준-맵핑 참조.
2. NIA 채은선, "인공지능기본법의 주요 내용" 세미나(AI 제품안전 혁신 TF 제6차 전체 회의, 2026-08-06). 고영향 AI 판단 가이드라인(법 제33조③)은 NIA, 고영향 AI 사업자 책무 고시·가이드라인(법 제34조)은 TTA가 작성 주체다. 상세는 AI기본법-세미나-260806-NIA채은선-요약·한국AI기본법-EU-AI-Act-조문대응 § 13-1 참조.

15장. 규제를 사업 기회로 — 조달·표준·인증 시장

이 장을 한 문장으로

규제는 비용이기만 한 것이 아니다. 남들이 넘지 못한 문턱은 이미 넘은 쪽에게 울타리가 된다.

운전면허를 생각해 보자.

면허를 따는 일은 순전히 비용이다. 학원비가 들고 시간이 들고 시험에 떨어지면 다시 봐야 한다. 면허를 판다고 차가 더 잘 나가지도 않는다.

그런데 면허 제도에는 다른 면이 있다. **면허가 없는 사람은 도로에 나올 수 없다.** 이미 면허를 가진 사람에게 이 제도는 비용인 동시에, 준비되지 않은 운전자를 도로에서 배제해 주는 장치다.

한 가지가 더 있다. 면허 제도는 그 자체로 산업을 만든다. 운전학원이 생기고, 교재가 팔리고, 시험을 대행하고 관리하는 일이 직업이 된다. **규제가 없었으면 존재하지 않았을 시장이다.**

AI 규제에서도 같은 구조가 만들어지고 있다. 다섯 갈래를 짚어 본다.

다섯 갈래

첫째, **표준을 만드는 자리.** 14장에서 본 조화표준은 회의실에서 만들어진다. 그 자리에 앉으면 자기 조직의 기술적 전제와 용어가 표준의

문안에 반영될 가능성이 생긴다. 표준이 확정된 뒤에 맞추는 쪽과, 만들어질 때 참여한 쪽은 같은 문서를 놓고도 준비 기간이 다르다.

둘째, ISO/IEC 42001 인증. 현재 인증받을 수 있는 유일한 AI 관리시스템 표준이다.^[1] 조직이 AI를 책임 있게 관리하는 체계를 갖췄음을 제3자가 증명해 준다.

여기서 정확히 해 둘 것이 있다. 이 인증은 두 법 어디에서도 준수 추정을 주지 않는다. 42001 인증서가 있다고 EU AI Act 제17조를 충족한 것이 되지는 않는다. 다만 그 요건에 대응하는 뼈대로 쓸 수 있고, 당국이나 거래 상대방에게 역량을 보여 주는 재료가 된다. **추정과 증명은 다르다.**

셋째, 우선구매 확인 취득. 7장에서 본 한국의 조달 지렛대다. 확인을 받으면 국가기관 우선구매 대상이 된다.

다만 실무적 기대치는 정확히 잡아야 한다. 확인 절차가 실제로 심사하는 것은 그 제품에 AI가 정말 들어 있는지이지 그 AI가 얼마나 안전한지가 아니다. 진입 비용이 낮다는 뜻이고, 따라서 이 확인만으로 얻는 차별화 효과도 크지 않다는 뜻이다. **조달 시장 접근권이지 품질 신호가 아니다.**

넷째, 역외 대리인 역할. EU AI Act는 역내에 거점이 없는 공급자에게 수권대리인 지정을 요구한다. 한국 법도 일정 기준을 넘는 해외 사업자에게 국내대리인 지정을 요구한다(11장).

양쪽 모두 누군가는 그 역할을 맡아야 한다. 법률·기술 문서를 이해하고 당국과 소통할 수 있는 조직에게 이것은 서비스 상품이 된다.

다섯째, 적합성평가 시장. 고위험 AI 중 일부는 제3자 검증기관의 평가를 거쳐야 한다. 2026년 개정은 검증기관 지정 범위를 코드 체계로 세분화했다. 제품 관련 코드, 생체인식 코드, 그리고 기술 유형별 코드가

나눠는데 마지막 항목에는 에이전틱 AI를 포함한 신흥 기술까지 들어가 있다.^[2]

평가를 받아야 하는 쪽이 있으면 평가를 하는 쪽도 있어야 한다. 이 시장은 이제 열리는 중이다.

[필자 분석] 누구에게 기회인가

지금까지의 이야기는 "규제를 비용이 아니라 기회로 보라"는 흔한 조언과 닮았다. 그런데 이 조언에는 잘 언급되지 않는 전제가 있다.

규제를 기회로 바꾸는 일에도 자원이 든다.

표준 회의에 사람을 계속 보내려면 그 사람의 시간을 살 수 있어야 한다. 인증을 받으려면 심사 비용과 준비 인력이 필요하다. 대리인 서비스를 하려면 법률과 기술을 함께 아는 인력을 데리고 있어야 한다.

이 모든 것이 이미 자원을 가진 쪽에 유리하다. 규제가 만든 새 시장은 규제를 감당할 수 있었던 조직들이 나눠 갖는다. 문턱이 울타리가 된다는 말은, 뒤집으면 울타리 바깥에 남는 쪽이 생긴다는 말이다.

EU가 이 문제를 모르는 것은 아니다. 2026년 개정에서 중소기업 지원 조항을 넓히면서 SME를 갓 졸업한 중견기업(SMC)이라는 범주를 새로 만든 것이 그 증거다. 기술문서를 간소화된 양식으로 낼 수 있게 하고, 품질관리시스템 요건을 완화하고, 규제 샌드박스에 우선 접근권을 준다. 규모가 커졌다는 이유만으로 갑자기 대기업과 같은 부담을 지게 되는 구간을 완충하려는 설계다.

그러니 "규제는 기회다"라는 문장은 절반만 맞다. 나머지 절반은 이것이다. 누구에게 기회인지를 빼고 말하면, 그 문장은 이미 앞서 있는 쪽의 언어가 된다.

여기까지가 기업의 눈으로 본 두 법이다. 남은 것은 하나, 이 모든 차이를 하나의 저울에 올리는 일이다.

실무자를 위한 체크박스

- ☐ 우리 분야의 조화표준이 어디서 누구에 의해 만들어지고 있는지 파악했는가 — 참여 여부는 그다음 문제다
- ☐ ISO/IEC 42001 인증을 검토한다면, 그것이 **준수 추정**이 아니라 **증명 재료**임을 조직 내에서 정확히 공유하고 있는가
- ☐ 우선구매 확인을 준비한다면, 그 효과를 조달 접근권으로 한정해 기대치를 설정했는가
- ☐ EU 진출 시 수권대리인을 직접 둘 것인지 외부 서비스를 쓸 것인지 결정했는가
- ☐ 우리 조직이 이 시장에서 **평가받는** 쪽인지 **평가하는** 쪽이 될 수 있는지 검토해 봤는가

한줄 코멘트.

규제를 기회로 바꾸는 일에도 자원이 든다. 그 말을 빼면 '기회'라는 단어는 이미 앞선 쪽의 언어가 된다.

주(註)

1. ISO/IEC 42001:2023 (Information technology — Artificial intelligence — Management system). 인증 가능한 AI 관리시스템 표준으로는 현재 유일하다. EU AI Act Article 17 품질관리시스템과 조항 대응 관계가 있으나 **자동 준수를 보장하지 않는다**. 상세는 AI관리시스템 참조.
2. Regulation (EU) 2026/1744 Article 1(43), AI Act Annex XIV 신설. 검증기관 지정 범위를 세 갈래 코드로 특정한다 — AIP(Annex I 제품 관련), AIB(Annex III 제1호 생체인식), AIH(기술 유형별). AIH 0401은 "다른 코드에 포함되지 않는 신흥 AI 기술(에이전틱 AI 포함)"이다.

에필로그 — 어느 쪽이 더 똑똑한 규제인가

처음에 던진 질문으로 돌아갈 차례다.

같은 기술을 향해 두 개의 법이 거의 동시에 작동을 시작했다. 열다섯 장에 걸쳐 그 둘을 조문 단위로 나란히 놓았다. 그러면 이제 답할 수 있어야 한다. 어느 쪽이 더 똑똑한 규제인가.

채점표를 펴 보자.

일곱 항목의 채점

1. **금지 목록** — EU는 열 가지를 아예 만들지 못하게 했고, 한국에는 그 목록 자체가 없음.
2. **분류 체계** — EU는 제품에 붙은 AI와 독립 소프트웨어를 나눠 각각 다른 경로로 보내고, 한국은 열 개 영역을 한 줄로 나열함.
3. **표시 의무** — EU는 기계가 읽을 수 있는 형식으로 새기라 하고, 한국은 사람이 읽는 안내 문구로도 갈음할 수 있게 함.
4. **기술 요건** — EU는 데이터 거버넌스·기술문서·자동 로그·정확성·견고성까지 여덟 갈래를 요구하고, 한국은 여섯 가지 책무를 두되 그중 절반가량이 EU에 대응물이 없음.
5. **사전 관문** — EU는 적합성평가를 통과해야 시장에 낼 수 있고, 한국은 검·인증을 노력 의무로 두고 조달로 유도함.
6. **사후 체계** — EU는 중대사고 보고와 사후 시장 모니터링을 사업자 의무로 부과하고, 한국에는 대응 조문이 없음.

7. 제재 — EU는 전 세계 매출의 7%까지 물리고, 한국의 과태료 상한은 3천만원임.

일곱 항목 전부 EU가 앞선다. 채점은 끝난 것처럼 보인다.

채점표를 누가 만들었는가

그런데 이 채점표에는 문제가 하나 있다.

항목을 EU의 조문 목록에서 뽑았다.

데이터 거버넌스가 있는가, 자동 로그가 있는가, 적합성평가가 있는가. 이 질문들은 EU AI Act의 목차를 그대로 옮긴 것이다. 한 법의 설계를 체크리스트로 만들어 다른 법을 채점하면, 다른 법은 반드시 진다. 애초에 그 설계대로 지어지지 않았기 때문이다.

0-3에서 이 책의 독법을 선언하며 말했다. 같은 법을 읽고도 다른 결론이 나오는 이유는 읽는 방법이 다르기 때문이라고. 채점에도 같은 문제가 있다. 무엇을 항목으로 세울지 정하는 순간 승부는 절반쯤 결정된다.

그렇다면 한국의 설계도로 채점표를 만들면 어떻게 되는가. 산업 육성 조문이 몇 개인가, 취약계층 지원 체계가 있는가, 학습용 데이터 제공 시스템이 있는가, 창업 투자 연계가 있는가. 이 항목들로 재면 EU AI Act는 거의 모든 칸이 비어 있다. EU가 규제 법령이기 때문이다.

규제 쪽에서도 마찬가지다. 국가가 스스로 구매자가 되어 조달 조건으로 검·인증을 요구하는 방식(7장)은 EU AI Act에 대응 조문이 없다. 강제하는 대신 유인하는 이 설계를 유럽식 체크리스트로 재면 "적합성평가 없음"이라는 빈칸으로만 잡힌다. 있는 것이 없는 것으로 기록되는 것이다.

두 법은 서로 다른 질문에서 태어났다(1장). 그 사실을 확인해 놓고 한쪽의 답안지로 다른 쪽을 채점하는 것은 앞뒤가 맞지 않는다.

강한 쪽이 늦게 온다

채점표가 놓치는 것이 하나 더 있다. 시점.

일곱 항목에서 앞선 EU의 고위험 AI 요건은 2027년 12월에야 적용된다. 제품에 내장된 AI는 2028년 8월이다. 반면 한국 인공지능기본법은 2026년 7월에 이미 전면 시행됐다.

더 엄격한 법이 나중에 오고, 더 느슨한 법이 먼저 왔다.

그래서 지금 이 순간 실제로 지고 있는 의무의 무게는 채점표와 반대다. 한국에서 고영향 AI를 다루는 사업자는 이미 여섯 가지 책무 아래 있고, 유럽의 같은 사업자에게는 아직 1년 넘는 준비 기간이 남아 있다.

이 역전은 한시적이다. 2028년이 지나면 원래 순서로 돌아간다. 그러나 그때까지 몇 해 동안, 한국 기업은 국내 규제를 먼저 감당하면서 아직 오지도 않은 유럽 규제를 동시에 준비해야 한다. 규제 강도의 순위와 대응 순서의 순위가 어긋나는 구간이다.

그러면 무엇으로 재는가

질문을 바꿔야 한다. 어느 쪽이 더 똑똑한가가 아니라, 어느 쪽이 더 오래 버티는가.

기술은 계속 바뀐다. 법이 그때마다 개정을 따라가야 한다면 그 법은 언제나 늦는다. 40년을 버틴 유럽 체계가 그 문제를 푼 방식은 이랬다. 법은 결과만

정하고, 그 결과에 이르는 방법은 표준에 맡긴다. 표준은 법보다 빨리 고칠 수 있다. 그래서 기술이 앞서 나가도 법의 골격은 흔들리지 않는다.

이 구조가 있느냐 없느냐가 내구성을 가른다. 그리고 이것은 EU 조문을 몇 개 베껴 온다고 생기지 않는다.

조문의 수가 아니라 조문 뒤에 서 있는 것의 두께다.

11장에서 두 나라의 감시 조직을 놓고 물었던 질문이 여기서 다시 나온다. 한국은 한 부처가 총괄하니 EU가 애써 폰 조정 문제 자체가 없다. 그것은 효율일 수도 있고, 아직 아무것도 짓지 않아서 생긴 빈칸일 수도 있다. 둘을 가르는 것은 조직도가 아니라 그 아래 무엇이 깔려 있는가였다.

같은 기준으로 다시 보면 한국의 공백은 다섯 개다. 금지 목록이 없고, 사전 관문이 노력 의무이며, 준수를 추정해 줄 표준 체계가 없고, 사고를 모아 되먹이는 회로가 없고, 어졌을 때의 대가가 가볍다.

이 다섯은 서로 다른 다섯 개의 결함이 아니다. 전부 하나의 사실에서 나온다. 위에 엮을 40년치 기초가 아직 없다는 것.

EU도 완주하지는 못했다

그렇다고 유럽이 답을 다 가진 것은 아니다.

AI 피해자의 입증 부담을 덜어 주려던 별도 지침은 결국 철회됐다(10장). 남은 것은 제조물책임 쪽으로 흡수된 부분적 장치이고, 그마저 제품에 한정된다. 순수한 서비스형 AI가 사람을 잘못 판단했을 때의 구제는 유럽에서도 여전히 회색지대다.

조화표준도 아직 도착하지 않았다. 품질관리시스템에 관한 첫 표준은 이 책을 쓰는 시점까지 초안 단계에 머물러 있다(14장). 준수 추정이라는 유럽

체계의 핵심 장치가 정작 AI 영역에서는 아직 작동하지 않고 있는 것이다.

"EU가 앞서 있다"와 "EU가 해결했다"는 다른 말이다. 이 구분을 놓치면 한국의 과제를 베끼기로 오해하게 된다.

40년을 누가 만드는가

유럽의 40년은 저절로 쌓이지 않았다. 사고가 났고, 소송이 걸렸고, 국가끼리 다투고, 그때마다 조금씩 고쳤다. 표준을 만드는 회의실에 사람들이 모여 문장 하나를 두고 몇 달을 보냈다.

두꺼란 그렇게 만들어진다. 그리고 그 자리에 **앉을 수 있는 사람과 앉을 수 없는 사람이 나뉜다**(15장). 규제를 감당할 자원이 있는 쪽이 규제의 문장을 쓰고, 그 문장이 다시 나머지의 조건이 된다.

이 책의 제목에 쓴 '설계자들'은 입법자만을 가리키지 않는다. 표준을 쓰는 사람, 인증을 하는 사람, 그리고 그 기준에 맞춰 물건을 만드는 사람 전부가 설계에 참여하고 있다. 다만 참여의 크기가 같지 않을 뿐이다.

프롤로그에서 이렇게 적었다. 우리는 얼굴을 모르는 사람이 만든 물건에 매일 목숨을 맡기고 있으며, 이제 맡기는 것은 물건이 아니라 판단이라고.

그 판단을 믿을 수 있게 만드는 일이 규제다. 규제는 기술을 묶는 밧줄이 아니라 **신뢰를 지탱하는 구조물**이다. 신뢰가 없으면 시장이 서지 않고, 시장이 서지 않으면 기술도 자라지 않는다.

그래서 마지막 질문은 어느 쪽이 이겼는가가 아니다.

지금 만들어지고 있는 이 구조물이 40년 뒤에도 서 있을 것인가. 그리고 그때 그 구조물을 설계한 사람들의 명단에, 누구의 이름이 올라가 있을 것인가.

후기 — 규제 지식베이스는 어떻게 만드는가

14장에서 어디를 봐야 하는지를 다뤘다. 이 후기는 그다음 문제를 다룬다. **본 것을 어디에 쌓을 것인가.**

규제 담당자의 자료는 대개 두 곳에 흩어진다. 받은 메일함과 개인 폴더. 둘 다 검색은 되지만 연결은 되지 않는다. 반년 뒤 "이 조문이 왜 이렇게 바뀌었더라"를 물으면 답이 나오지 않는다.

아래는 조직 규모와 무관하게 재현할 수 있는 방식이다.

세 층으로 나눈다

1층, 원문. 법령·고시·표준·가이드라인의 원본 파일을 그대로 모아 두는 곳이다. 요약하지 않고, 고치지 않는다. 관보 PDF면 관보 PDF 그대로 둔다. **이 층은 읽기 전용이다.**

2층, 정리. 원문 하나마다 요약 문서를 하나 만든다. 무엇을 규정하는지, 어떤 조문이 핵심인지, 다른 법령과 어떻게 연결되는지. 개념 단위 문서(예: '조화표준', '적합성평가')를 따로 두고 요약 문서와 서로 링크한다.

3층, 분석. 두 개 이상의 자료를 엮어야 나오는 것들이다. 조문 대응표, 일정 비교, 공백 진단. 이 책의 부록이 그 예다.

핵심은 **3층이 2층을, 2층이 1층을 가리키게 하는 것이다.** 어떤 주장이든 링크를 따라 내려가면 원문에 닿아야 한다. 닿지 않는 주장은 근거가 없는 주장이다.

어디서부터 시작하는가

한 번에 전부 만들려 하면 시작하지 못한다.

법령 하나를 골라 3층까지 끝까지 만들어 본다. 우리 사업에 가장 직접 걸리는 것 하나면 된다. 원문을 넣고, 요약을 쓰고, 다른 자료와 연결해 본다. 이 한 바퀴를 돌고 나면 이후 작업의 형식이 정해진다.

두 번째 법령부터 연결이 생긴다. 첫 문서는 링크할 곳이 없지만 두 번째부터는 있다. 그리고 세 번째쯤에서 두 자료가 서로 어긋나는 지점이 보이기 시작한다. **그 어긋남이 분석의 재료다.** 조문 대응표든 공백 진단이든, 전부 어긋남을 발견한 자리에서 나온다.

빈칸을 미리 만들어 두지 않는다. 언젠가 쓸 것 같은 폴더와 서식을 먼저 만들어 두면 대부분 빈 채로 남는다. 필요해진 순간에 만드는 편이 낫다.

지킬 것 네 가지

첫째, 1차 소스를 이긴 2차 자료는 없다. 보도와 해설은 원문을 찾아가는 안내판이지 원문의 대체물이 아니다. 조문의 내용을 단정하려면 조문을 읽어야 한다. 요약본만 보고 "이 법은 이렇게 규정한다"고 적으면, 요약한 사람의 실수가 그대로 옮겨 온다.

둘째, 확인하지 못한 것은 확인하지 못했다고 적는다. 자료 안에 [검토 필요] 같은 고정 표기를 하나 정해 두고, 확인에 실패한 지점마다 붙인다. 빈칸으로 두면 나중에 읽는 사람이 그것을 확인된 사실로 오해한다.

셋째, "아직 없다"에는 날짜를 붙인다. 14장에서 다룬 함정이다. 없다는 사실은 있다는 사실보다 빨리 낡는데, 낡아도 티가 나지 않는다. "고시 미제정"이라고만 적힌 문장은 고시가 나온 뒤에도 조용히 그 자리에

남는다. "2026년 8월 5일 확인 기준 미제정"이라고 적으면 유효기간이 함께 적힌다.

넷째, 정정은 지우지 말고 남긴다. 틀린 것을 발견하면 조용히 고치고 싶어진다. 그러나 무엇이 왜 틀렸는지를 함께 남겨야 같은 실수가 반복되지 않는다. 특히 원문을 대조하지 않고 "저 자료가 틀렸다"고 단정하는 일은 생각보다 자주 생긴다. 그런 판단이야말로 기록으로 남겨 두고 검증받아야 한다.

생성형 AI를 쓸 때

이 작업에 생성형 AI를 도구로 쓸 수 있다. 원문을 읽히고 요약하게 하고 대조표를 만들게 하는 일은 잘 맞는다. 다만 제약이 분명하다.

모델의 기억은 원문이 아니다. 학습된 지식으로 조문을 말하게 두면 그럴듯하지만 틀린 조문 번호가 나온다. 반드시 원문 파일을 함께 주고, 그 파일 안에서만 답하게 한다.

부정형 판단을 맡기지 않는다. "이런 조문은 없다"는 결론은 전수 확인을 거쳐야 나올 수 있는데, 모델은 찾지 못한 것과 없는 것을 구분하지 못한다.

작업 규칙을 문서로 만들어 함께 준다. 어떤 표기를 쓸지, 무엇을 확인해야 하는지, 어떤 구역은 손대면 안 되는지. 규칙이 파일로 있으면 사람이 바뀌어도 결과가 흔들리지 않는다.

누가 갱신하는가

가장 흔한 실패는 자료가 없어서 생기지 않는다. 자료가 낡아서 생긴다.

정리해 둔 문서에는 만든 날짜가 아니라 **마지막으로 원문과 대조한 날짜**를 적는다. 그리고 무엇을 얼마나 자주 다시 볼지를 자료 안에 적어 둔다. 매달 확인할 것, 분기에 한 번 볼 것, 확정되면 그때 손댈 것.

이것이 정해져 있지 않으면 지식베이스는 어느 한 시점의 사진으로 굳는다. **굳은 자료는 없는 자료보다 위험하다.** 읽는 사람이 그것을 현재로 착각하기 때문이다.

마지막으로

이 방식의 목적은 자료를 많이 모으는 것이 아니다.

반년 뒤의 자신이 오늘의 판단 근거를 되짚을 수 있게 하는 것이다. 규제는 계속 바뀌고, 바뀔 때마다 우리는 "그때 왜 그렇게 판단했지"를 묻게 된다. 그 질문에 답할 수 있는 조직과 답할 수 없는 조직의 차이는, 자료의 양이 아니라 자료가 서로 연결되어 있는지에서 갈린다.

APPENDIX

부록

조문 대응, 시행 일정, 용어, 원문 출처.

부록 A. 한국 인공지능기본법 ↔ EU AI Act 조문 대응표

본문에서 다른 대응 관계를 한 표로 모았다. 왼쪽 열의 장을 먼저 읽고, 이 표는 나중에 조문 번호를 찾을 때 펼치는 색인으로 쓴다.

상태 표기 —  **확정:** 법령·시행령·고시가 모두 나와 있다.  **고시 대기:** 조문은 확정됐지만 그 조문을 실행할 고시가 아직 없다. 상태는 2026-08-05 기준이며, 고시는 수시로 새로 나오므로 실무 적용 전에는 다시 확인한다(14장·부록 B).

축	한국 조문	EU 조문	핵심 차이	상태
정의	법 제2조1·2호	Art.3(1)	한국은 지적능력 구현 중심, EU는 자율성·적응성·추론 중심(OECD 정의 계열)	
리스크 분류	법 제2조4호(고영향 10영역)	Art.6+Annex III(8카테고리)+Annex I(제품 내장)	한국은 단일 목록, EU는 "제품 부품형/독립 시스템형" 이원 체계 (3장)	
금지 AI	대응 조문 없음	Art.5(1)(a)-(j), 10개	한국은 금지 카테고리 자체가 없다 — 전부 "고영향"으로 허용 후 책무 부과 (4장)	

축	한국 조문	EU 조문	핵심 차이	상태
제품 내장 AI 분류	별표 1(5개 분야 이행간주)	Annex I Section A(11개)·Section B(기계류 등 9개, 2026 개정으로 편입)	EU는 조화법령 소속에 따라 적용 강도가 갈린다 (5장)	✅
표시·워터마크	법 제31조1-3항	Art.50(1)(2)(4)	한국은 기계판독 표시가 선택, EU는 의무 (6장)	✅ (경과조치 2026.12.02까지)
검·인증·우선구매	법 제30조, 시행령 제15조④1호·⑤	Art.40(조화표준 준수추진)·Art.43(강제 적합성평가)	한국은 조달이라는 한 구역에서만 지렛대로 작동, EU는 시장 전체의 관문 (7장)	✅ 확인 고시 2026-07-21 제정
영향평가	법 제35조(노력 의무)	Art.27 FRIA(배포자 의무)	같은 방향이나 강도가 다르다 — 노력 대 의무	✅
프론티어·GPAI	법 제32조, 시행령 제24조	Art.51-55	연산량 문턱 한국 10 ²⁶ FLOPs, EU 10 ²⁵ FLOPs — 한국이 10배 관대 (8장)	⚠️ 연산량 산정 방식 세부 고시 대기
사업자 6책무	법 제34조	Art.8-15(8대 기술요건)	한국은 6개 책무 나열, EU는 데이터 거버넌스·자동 로그·정확성·견고성까지 요구	✅ (한국 쪽 세부는 고시 대기 항목 있음)
감시·집행	법 제40조(사실 조사)	Reg 2019/1020 + AI Act Ch.IX, 2026 개정으로 AI Office 권한 대폭 확장(신설 Art.75a-75d)	한국은 행정조사 수준, EU는 위장 구매·온라인 차단명령까지 포함한 법정 최소 권한 (11장)	✅

축	한국 조문	EU 조문	핵심 차이	상태
국내대리인	법 제36조, 시행령 제29조	Art.22	한국은 매출·이용자 수 기준 선별 규제, EU는 역외 공급자 전원에 무차별 적용	✓
제재 상한	법 제43조(과태료 최대 3천만원)	Art.99-101(매출 7% 또는 €35M)	대기업 기준 수백~수천 배 격차 (9장)	✓
시행 일정	법 2026.07.21 전면 시행	Annex III 2027.12.02 / Annex I 2028.08.02(2026 확정 개정)	한국이 EU보다 먼저 규제 적용을 받는 역전 구간 발생	✓

고시 위임 항목 중 의무·미제정 3종 — 연산량 산정방식(시행령 제24조②), 프론티어 이행방식·결과제출(법 제32조 ③), 고영향 영향평가 세부 기준(시행령 제28조 ③). 재량 개방 2종(투명성 예외 추가 사유·고영향 확인 고려사항 추가)은 미제정이어도 제도 작동에는 지장이 없다. 세 구분과 상세는 부록 B에서 다룬다.

관련 문서

- 한국AI기본법-EU-AI-Act-조문대응 — 이 표의 원 근거, 12절 상세 분석
- Digital-Omnibus-Omnibus-VII — EU 쪽 2026년 개정 상세
- 인공지능제품서비스확인절차고시 — 검·인증·우선구매 행의 근거

부록 B. 시행 일정 캘린더 2026~2030

14장(정책 환경을 읽는 법)이 참조하는 실무 도구다. 왼쪽에 날짜, 오른쪽에 그날 무엇이 시작되는지를 적었다.

일정표

시점	한국	EU
2026.07.21	인공지능기본법 전면 시행 (법·시행령)	—
2026.07.21	우선구매 확인 고시 제정·시행	—
2026.07.27	—	Digital Omnibus on AI 발효 (Regulation (EU) 2026/1744)
2026.08.02	—	Art.50 일반 투명성 의무 적용
2026.12.02	—	워터마킹 기계판독 표시 경과 조치 종료 (2026.08.02 이전 출시분)
2026.12.02	—	CSAM·비동의 친밀 이미지 (NCII) 생성 AI 금지 적용
2027.08.02	—	회원국 AI 규제 샌드박스 구축 마감
2027.12.02	—	Annex III 독립형 고위험 AI 전면 적용
2028.01.28	—	검증기관(Notified Body) 경과 조치 종료

시점	한국	EU
2028.08.02	—	Annex I 제품 내장형 고위험 AI 전면 적용, 기계류규정 흡수 위임입법 적용 시한
2030.08.02	—	공공기관이 배치하는 고위험 AI 전면 적용

한국 하위 고시 대기 목록 (2026-08-05 확인 기준)

법령·시행령은 확정됐으나 그것을 실행할 고시가 아직 없는 항목이다.

빠진 고시가 다 같은 무게는 아니다. 조문 문언을 보면 두 종류가 섞여 있다. 반드시 만들어야 하는 것(고시한다·고시하여야 한다)과, 필요하면 추가할 수 있게 열어 둔 것(고시하는 경우·고시하는 사항)이다. 앞의 것이 없으면 제도가 작동하지 않고, 뒤의 것이 없어도 제도는 돌아간다.

A. 만들어야 하는데 아직 없는 것

항목	위임 조문	문언	상태
프론티어 AI 연산량 산정 방식	시행령 제24조②	고시한다	⚠ 대기
프론티어 AI 이행 방식·결과 제출 방법	법 제32조③	고시하여야 한다	⚠ 대기
고영향 AI 영향평가 그 밖에 필요한 사항	시행령 제28조③	고시한다	⚠ 대기

B. 필요할 때 추가하도록 열어 둔 것

항목	위임 조문	문언	상태
투명성 고지·표시 적용 예외의 추가 사유	시행령 제23조④제3호	고시하는 경우	미제정(제도 작동에 지장 없음)
고영향 AI 확인 시 고려할 추가 사항	시행령 제25조②제5호	고시하는 사항	미제정(제도 작동에 지장 없음)

C. 이미 나온 것

항목	위임 조문	상태
우선구매 확인 기준·절차	시행령 제15조④제1호·⑤	 제정 — 「인공지능제품·서비스 확인 절차 운영에 관한 고시」(과학기술정보통신부고시 제2026-50호, 2026.07.21)

가이드라인(비구속)은 고시(의무)보다 먼저 나올 수 있다. NIA 채은선 연구원의 2026.8.6. 세미나(AI 제품안전 혁신 TF)가 소개한 표에 따르면, A표의 세 항목 중 프론티어 AI 연산량 산정·이행방식은 이미 "AI 안전성 확보 가이드라인"이라는 이름의 비구속 가이드라인으로 판단기준·산정방식 상세가 나와 있다. 다만 가이드라인의 존재가 곧 시행령 제24조 ②·법 제32조③이 요구하는 의무 고시(관보 게재)의 제정을 뜻하지는 않는다 — 이 부록의 A표 상태는 그대로 유지한다. 같은 세미나에서 고영향 관련 고시·가이드라인의 작성주체가 NIA(판단 가이드라인)와 TTA(책무 고시·가이드라인)로 확인됐다.

"대기"의 범위를 오해하지 말 것. 위 항목들은 제도 전체가 비어 있다는 뜻이 아니다. 빠대는 이미 법과 시행령에 있고, 고시에 넘긴 것은 그 위에 없는 부분이다.

- **투명성 예외** — 시행령 제23조④가 이미 두 가지 예외를 직접 정하고 있다. AI를 썼다는 사실이 제품명·화면 문구 등으로 **명백한 경우**,

그리고 사업자의 내부 업무 용도로만 쓰이는 경우다. 고시는 그 밖의 예외를 추가할 때 필요하다.

- **영향평가** — 시행령 제28조 ① 이 영향평가에 반드시 답아야 할 일곱 항목을 이미 열거하고 있다(영향받을 개인·집단 식별, 관련 기본권 유형, 사회적·경제적 영향, 사용 행태, 평가지표와 산출 방식, 위험의 예방·완화·복구, 개선 이행계획). 고시는 그 외 필요사항을 정한다.
- **고영향 확인 절차** — 시행령 제25조가 제출 서류, 판단 고려사항, 회신 기한 30일(1회 30일 연장 가능), 이의가 있을 때의 재확인 요청 10일 이내 절차까지 정해 두었다. 고시는 판단 시 고려할 사항을 추가할 때 필요하다.
- **프론티어 관련 두 항목** — 여기만은 대기의 무게가 다르다. 시행령 제24조 ①은 판정의 세 요건(누적 연산량 10^{26} 부동소수점연산 이상 · 최첨단 기술 적용 · 광범위하고 중대한 위험 우려)을 정했지만 그 연산량을 어떻게 세는지는 고시로 넘겼고, 법 제32조 ③은 안전성 확보 의무의 이행 방식과 결과 제출 방법을 장관이 고시하여야 한다고 못 박았다. 문턱은 세워 두고 채는 자를 아직 주지 않았으며, 지켜야 할 의무는 있는데 어떻게 지켰다고 인정받는지의 방법이 없다(8장).

시행일과 집행일은 다르다

표의 "2026.07.21 인공지능기본법 전면 시행"은 조문이 효력을 갖기 시작한 날일 뿐, 위반 시 실제로 조사·처분이 나가는 날은 아니다. 2026.8.6. 세미나 기준 사실조사와 과태료 부과는 유예 중이며, 인명사고나 심각한 인권 침해처럼 사회적으로 중대한 사안일 때만 예외적으로 사실조사에 들어간다(인공지능기본법 안내데스크 사례집, 2026, 7면). 이 표의 날짜를 읽을 때 "법이 켜졌다"와 "제제가 작동한다"를 같은 것으로 읽지 않는다.

어디를 보고 있어야 하는가

EU — 관보(Official Journal of the EU), AI Office 발표, 조화표준 진행 상황(CEN-CENELEC JTC 21), Safety Gate 공지.

한국 — 관보, 과학기술정보통신부 고시 행정예고, 국가인공지능전략위원회 의결 사항, 인공지능안전연구소 발표.

두 축 모두 이 책이 다루는 내용의 유효기간을 좌우한다. 이 표의 상태는 고정된 사실이 아니라 특정 시점의 스냅샷이다 — 실무에 쓸 때는 반드시 다시 확인한다.

확인 방법. 한국 고시는 국가법령정보센터의 **행정규칙** 검색으로 확인한다. 위 대기 목록의 상태는 2026년 8월 5일에 행정규칙을 직접 조회해 확인했다(「고영향」·「연산량」·「생성형」 검색 결과 0건, 「인공지능」 검색 결과 중 과학기술정보통신부 고시는 제2026-50호 하나). 법령 본문은 시행령 제23·24·25·28조를 조문 단위로 대조했다.

관련 문서

- Digital-Omnibus-Omnibus-VII — EU 쪽 일정의 원 근거
- 한국-AI기본법-시행령36506호 · 인공지능제품서비스확인절차고시 — 한국 쪽 근거
- book/appendix/A-조문대응표 — 조문별 상세 대응
- AI기본법-세미나-260806-NIA채은선-요약 — 고시·가이드라인 작성주체, 집행유예 상태 출처

부록 C. 용어 사전

본문에서 처음 나올 때 괄호로 풀어 둔 용어를 한자리에 모았다. 괄호 안 숫자는 그 용어를 자세히 다룬 장이다.

ㄱ~ㄴ

가지번호 — 법을 고치면서 기존 조문 사이에 새 조문을 끼워 넣을 때 붙이는 번호. 제17조의2 같은 형태다. (1장)

검·인증등 — 한국 인공지능기본법 제30조가 만든 법률상 약칭. 검증과 인증을 묶어 부른다. (7장)

검증기관 — 유럽에서 국가가 지정해 적합성평가를 맡기는 제3자 기관. Notified Body. (0-2, 7장)

결함 추정 — 피해자가 결함을 직접 증명하지 못해도 일정 조건이 충족되면 결함이 있는 것으로 보는 법적 장치. (10장)

고영향 인공지능 — 한국법이 쓰는 분류. 열 개 영역에서 사람의 생명·신체 안전과 기본권에 영향을 줄 수 있는 AI. (3장)

고위험 AI — EU가 쓰는 분류. 제품에 내장된 안전부품형과 Annex III 독립형 두 갈래가 있다. (3장)

과징금과 과태료 — 과징금은 위반으로 얻은 이익을 환수하는 성격이 강하고, 과태료는 질서 위반에 대한 제재다. (9장)

금지 관행 — EU AI Act 제5조가 시장에 내놓는 것 자체를 막은 열 가지 AI. (4장)

기계판독 — 사람이 아니라 기계가 읽을 수 있는 형식으로 정보를 새기는 것. 워터마크나 메타데이터가 그 방법이다. (6장)

□~人

모듈 — 유럽 적합성평가의 표준 절차 묶음. A부터 H1까지 있고 제품의 위험도에 따라 어느 것을 쓸지 정해진다. (0-2)

배치자 — AI를 만들어 파는 쪽이 아니라 업무에 실제로 가져다 쓰는 쪽. EU가 공급자와 구분해 새로 세운 행위자다. (2장)

범용 AI 모델(GPAI) — 특정 용도에 묶이지 않고 여러 작업에 두루 쓰이는 모델. (8장)

부동소수점연산(FLOPs) — 소수점이 붙은 숫자를 더하고 곱하는 계산의 횟수. 모델 훈련에 들어간 계산의 총량을 재는 단위로 쓰인다. (8장)

수권대리인 — EU 안에 거점이 없는 공급자가 역내에 두어야 하는 법적 창구. 한국법의 국내대리인이 이에 대응한다. (11·15장)

시스템적 위험 — 모델이 충분히 크고 강력해서 사회 전반에 파급될 수 있는 위험. EU는 연산량 문턱으로 이를 추정한다. (8장)

실질적 변경 — 이미 시장에 나온 제품이 나중에 크게 바뀌어 새로 평가받아야 하는 상태. (5장)

실행규범 — 범용 AI 모델 공급자가 자율로 서명하는 문서. 서명하지 않으면 같은 의무를 혼자 입증해야 한다. Code of Practice. (7·8장)

○~ㅎ

안전부품 — 제품 안에서 안전 기능을 맡는 구성요소. 2026년 개정으로 정의가 좁아졌다. (12장)

위험원과 위해 — 위험원(Hazard)은 다치게 할 수 있는 잠재적 원천이고, 위해(Harm)는 실제로 생긴 피해다. 리스크는 이 둘을 조합해 잔다. (0-2)

적합성평가 — 제품이 법이 정한 요건을 충족하는지 확인하는 절차 전체. (0-2)

조화표준 — 유럽 표준화 기구가 만드는 기술 문서. 이것을 지키면 법 요건을 충족한 것으로 추정해 준다. (0-2, 14장)

준수 추정 — 조화표준을 따랐으면 법 요건도 충족한 것으로 보아 주는 장치. 유럽 체계의 핵심 부품이다. (0-2)

품질관리시스템(QMS) — 문서 한 벌이 아니라 그 문서를 계속 갱신하고 유지하는 조직 체계. (13장)

프론티어 AI — 한국법이 쓰는 표현. 최첨단 기술로 만들어진 대형 모델을 가리킨다. (8장)

행정예고 — 고시를 확정하기 전에 내용을 미리 알리고 의견을 받는 절차. 개입할 수 있는 지점이다. (14장)

알파벳

Annex I / Annex III — EU AI Act의 두 부속서. 앞은 제품안전법 목록이고 뒤는 독립형 고위험 AI 여덟 영역이다. (3·12장)

CE 마킹 — 제조사가 유럽 법 요건을 충족했다고 스스로 선언하며 붙이는 표시. 품질 보증서가 아니다. (0-2)

ISO/IEC 42001 — AI 관리시스템 국제표준. 현재 인증받을 수 있는 유일한 AI 관리시스템 표준이다. (15장)

NLF — New Legislative Framework. 1985년의 New Approach, 1989년의 모듈, 2008년의 인정·시장감시를 하나로 묶은 유럽 제품안전의 공통 문법. (0-1, 2장)

prEN / EN — 유럽표준의 단계 표시. **pr** 이 붙어 있으면 아직 제안 단계이고, 떨어지면 정식 표준이다. (14장)

Safety Gate — 위험한 제품 정보를 회원국끼리 즉시 돌려 보는 EU의 신속경보 체계. 옛 이름은 RAPEX다. (0-2, 14장)

SME / SMC — 중소기업, 그리고 중소기업을 갓 졸업한 중견기업. 2026년 개정이 후자를 새 범주로 만들었다. (15장)

부록 D. 원문·출처 목록

이 책의 주장은 전부 아래 원문에 닿아야 한다. 요약본이나 해설이 아니라 관보에 실린 문서를 기준으로 삼았다.

두 AI법과 하위법령

- 인공지능 발전과 신뢰 기반 조성 등에 관한 기본법 (법률 제21311호, 2026-07-21 시행)
- 같은 법 시행령 (대통령령 제36053호 → 제36506호 일부개정, 2026-07-21 시행) 및 별표 1·2
- 인공지능제품·서비스 확인 절차 운영에 관한 고시 (과학기술정보통신부고시 제2026-50호, 2026-07-21)
- Regulation (EU) 2024/1689 — Artificial Intelligence Act
- Regulation (EU) 2026/1744 — Digital Omnibus on AI (AI Act 개정, 2026-07-27 발효)

EU 제품안전의 골격

- Decision No 768/2008/EC — 적합성평가 모듈과 공통 문법
- Regulation (EC) No 765/2008 — 인정 및 CE 마킹 일반원칙
- Regulation (EU) No 1025/2012 — 유럽 표준화, 준수 추정근거
- Regulation (EU) 2019/1020 — 시장감시
- The 'Blue Guide' on the implementation of EU product rules 2022 (OJ C 247) — 법령이 아닌 집행위 해설

섹터 법령

- Regulation (EU) 2023/1230 (기계류) · Regulation (EU) 2017/745 (의료기기) · Regulation (EU) 2017/746 (체외진단) · Directive 2014/53/EU (무선기기) · Directive 2014/35/EU (저전압) · Directive 2009/48/EC 및 Regulation (EU) 2025/2509 (완구)

책임과 안전망

- Directive (EU) 2024/2853 — 신 제조물책임지침
- Council Directive 85/374/EEC — 구 제조물책임지침
- Regulation (EU) 2023/988 — 일반제품안전규정
- Regulation (EU) 2024/2847 — 사이버복원력법
- Regulation (EU) 2022/2065 — 디지털서비스법

표준·가이드라인

- ISO/IEC Guide 51 (안전 측면) · ISO/IEC 42001 (AI 관리시스템) · ISO/IEC 23894 (AI 리스크관리) · ISO/IEC 22989 (AI 개념과 용어)
- prEN 18286 — AI Act용 품질관리시스템 유럽표준 초안 (2026-08-05 기준 제안 단계)
- GPAI Code of Practice — 투명성·저작권·안전보안 3개 챕터

어디서 찾는가

한국 법령·고시는 국가법령정보센터에서, EU 법령은 EUR-Lex에서 전문을 볼 수 있다. 두 곳 모두 조문별 시행일과 개정 이력을 함께 제공한다. 인용할

때는 법령명과 조문 번호에 더해 시행일 또는 조회일을 함께 적는다 —
같은 조문이라도 시점에 따라 내용이 다르기 때문이다.