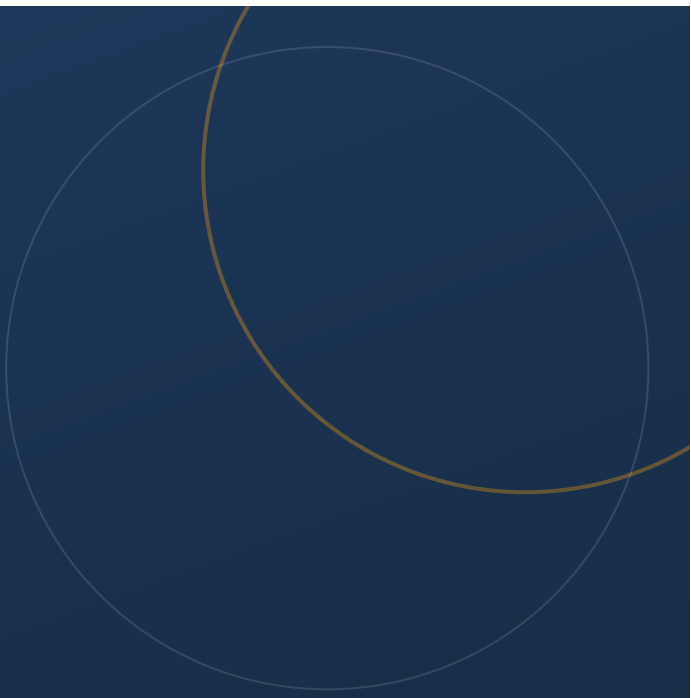




POROZUMIEŤ EURÓPSKEJ REGULÁCII AI

EU AI Act: úvod

Od zakázaných praktík po povinnosti pre vysokorizikové systémy:
sprievodca európskou reguláciou AI v zrozumiteľnom jazyku

Robert Barcik

Prvé vydanie, júl 2026 · revízia a slovenské vydanie august 2026

EU AI Act: úvod · Prvé vydanie, júl 2026 · revízia a slovenské vydanie august 2026

© 2026 Robert Barcik, LearningDoe s.r.o.. Všetky práva vyhradené.

Táto kniha je sprievodnou učebnicou k online kurzu „Úvod do EU AI Act“.

Táto kniha je vzdelávací materiál, nie právne poradenstvo. AI Act EÚ a jeho usmernenia sa ďalej vyvíjajú; pri rozhodnutiach s právnym dosahom sa poraďte so svojím právnym zástupcom.

Slovenské vydanie preložil Claude (Fable 5) 16. augusta 2026, významovo, nie slovo za slovom; citácie zákona sú prevzaté z úradného slovenského znenia nariadenia (EÚ) 2024/1689 v konsolidovanom znení k 27. júlu 2026. Rámčeky „Stav k...“ zodpovedajú augustovej revízii anglického originálu. Pri pochybnostiach platí [anglický originál](#).

Úvod: ako čítať túto knihu

Niekde v Európe práve teraz banka rieši, či sa jej model na schvaľovanie úverov počíta ako umelá inteligencia. Marketingový tím sa pýta, či chatbot na ich webe musí oznámiť, že je chatbot. Personálnej riaditeľke ponúkajú nástroj, ktorý zoraďuje uchádzačov o prácu, a niečo na ňom ju znepokojuje, hoci zatiaľ nevie povedať čo.

EU AI Act je zákon, ktorý odpovedá všetkým trom. Táto kniha je sprevádzaná prehliadka po ňom.

Ešte pred všetkým ostatným tu je upozornenie, ktoré dávam na začiatku každého kurzu, každej prednášky a teraz aj každej knihy: **všetko tu berte s rezervou**. Som dátový vedec a lektor, nie váš právnik. Táto kniha je starostlivá rešerš preložená do zrozumiteľného jazyka a za tú starostlivosť si stojím, ale pri každom rozhodnutí so skutočným právnym dosahom sa poraďte so svojim právnym zástupcom. Za to, čo postavíte alebo rozhodnete, nemôžem niesť zodpovednosť. Tak, a teraz sa môžeme uvoľniť.

Čo táto kniha je

Toto je sprievodná učebnica k môjmu online kurzu *Úvod do EU AI Act*. Slúži rovnakému účelu a rovnakému publiku: ľuďom z biznisu, právnikom, produktovým manažérom, konzultantom a zvedavým profesionálom, ktorí potrebujú pracovné porozumenie európskej regulácii AI bez toho, aby sa brodili 180 odôvodneniami legislatívnej prózy.

Na čítanie knihy nepotrebuje kurz a na sledovanie kurzu nepotrebuje knihu. Pokrývajú tú istú látku v rovnakom poradí, takže môžete použiť ktorúkoľvek z nich samostatne, alebo knihu ako miesto, ku ktorému sa vrátite, keď si potrebujete overiť, čo vlastne hovorí článok 6 ods. 3. Videá sú lepšie v budovaní intuície krok za krokom. Kniha je lepšia v tom, že sa v nej dá hľadať, citovať z nej a čítať ju v posteli.

Čo tu nenájdete, je AI Act preusporiadaný a okomentovaný. Moja ambícia je iná: chcem, aby ste tejto regulácii *rozumeli* tak, ako rozumiete príbehu, aby ste vtedy, keď

vám na stôl pristane nová situácia, vedeli o nej uvažovať, namiesto toho, aby ste v miernej panike hľadali cez Ctrl+F v oficiálnom texte.

Čo budete vedieť na konci

Po poslednej kapitole by ste mali vedieť pozrieť sa na systém AI, skutočný alebo navrhovaný, a previesť ho otázkami, na ktorých záleží: Je to vôbec AI v zmysle Aktu? Je to, čo robí, priamo zakázané? Je vysokorizikový, a ak áno, čo to naozaj vyžaduje? Čo musí zverejniť ľuďom, ktorých sa dotýka? A kto v reťazci od vývojára po koncového používateľa nesie ktorú povinnosť?

Dúfam, že si odnesiete aj niečo vzácnejšie: kalibrovaný pocit, ako veľmi sa treba báť. Spoiler z polovice roka 2026: menej panicky, než naznačovali titulky, pozornejšie, než nerobiť nič.

Ako je kniha usporiadaná

Držíme sa štruktúry kurzu, ktorá sleduje logiku samotného Aktu. Kapitoly 1 a 2 stavajú základy: čo AI vlastne je a prečo sa Európa rozhodla, že potrebuje vlastný zákon. Kapitola 3 je celý Akt v kocke, vrátane termínov, pre čitateľov v časovej tiesni. Kapitoly 4 a 5 budujú vašu slovnú zásobu: definíciu systému AI a pojmy (poskytovateľ, nasadzujúci subjekt, profilovanie, zamýšľaný účel), na ktorých stojí všetko ostatné.

Potom rebrík rizika, zhora nadol. Kapitola 6 pokrýva zakázané praktiky, všetkých osem. Kapitoly 7 a 8 sa venujú vysokorizikovej AI, najprv tomu, ako systém skončí zaradený medzi vysokorizikové, potom tomu, aké povinnosti z toho plynú. Kapitola 9 pokrýva povinnosti v oblasti transparentnosti, kapitola 10 osobitný režim pre modely AI na všeobecné účely a kapitola 11 mašinériu okolo toho všetkého: regulátorov, experimentálne prostredia, sankcie a úprimný stav vymáhania. Kapitola 12 uzatvára praxou: čo naozaj urobiť v pondelok, plus bonusová prehliadka chaotického, fascinujúceho prieniku AI a autorského práva.

Každá kapitola končí krátkym kvízom. Nikto vás neznámkuje. Ale ak po káve dokázate odpovedať na päť otázok, kapitola splnila svoju úlohu.

Cestou stretnete aj štyri druhy rámciekov. Rámčeky **Akt hovorí** obsahujú vlastné slová zákona, pre chvíle, keď záleží na presnom znení. Rámčeky **Môj pohľad** som ja, keď vystúpim z roly sprievodcu, aby som vám povedal, čo si naozaj myslím. Rámčeky **Pozor** upozorňujú na nesprávne čítania, s ktorými sa stretávam najčastejšie. A rámčeky **Stav k** označujú fakty s dátumom spotreby: stav vymáhania, čakajúce usmernenia, čokoľvek, čo sa oplatí znovu overiť, kým sa na to spoľahnete. Všetko mimo rámciekov som stále ja, kto hovorí; rámčeky vám len hovoria, aký druh tvrdenia držíte v rukách.

Poznámka k dátumom, písaná v júli 2026 a revidovaná v auguste 2026

AI Act je živá vec. Účinnosť nadobudol v auguste 2024, jeho prvé povinnosti sa začali uplatňovať v roku 2025 a v roku 2026 už bol raz novelizovaný: takzvaný digitálny omnibus posunul viacero termínov, predovšetkým povinnosti pre vysokorizikové systémy na december 2027.

STAV K AUGUSTU 2026

Novela omnibus je zákon: nariadenie (EÚ) 2026/1744, uverejnené v Úradnom vestníku 24. júla 2026 a účinné od 27. júla 2026. Táto kniha všade používa dátumy po omnibuse, lebo práve tie riadia vaše plánovanie. Keď prvé vydanie v júli šlo do tlače, text ešte čakal na uverejnenie; augustová revízia aktualizovala rámčeky so stavom, ktoré cestou stretnete, a zvyšok knihy nechala tak, ako bol.

Tak, práve ste stretli svoj prvý rámček. Zhodou okolností nesie aj vašu prvú skutočnú lekciiu o regulácii EÚ, ešte než sa začne kapitola 1: dátumy sú zákon a zákon sa môže meniť. Kedykoľvek vám na nejakom termíne naozaj záleží, venujte dve minúty overeniu aktuálneho stavu, kým podľa neho konáte. Kniha vás naučí, kam sa pozeráť.

Nalejte si kávu. Začíname.

Ďalej: čo umelá inteligencia vlastne je, rozprávané cez šachový program, ktorý podvádzal memorovaním.

AI v zrozumiteľnej reči

Predstavte si, že ste v utorok ráno v pobočke banky. Chcete hypotéku. Pracovníčka sa pýta obvyklé otázky (váš príjem, vaše ostatné úvery, vaša pracovná situácia) a odpovede ťuká do počítača. Potom povie: „Dajte mi dvadsať sekúnd a poviem vám, či spĺňate podmienky.“

Dvadsať sekúnd. Nikto neodišiel do zadnej kancelárie. Nezasadala žiadna komisia. Niekde vnútri toho počítača sa kus softvéru pozrel na vaše čísla a rozhodol o jednom z najväčších finančných rozhodnutí vášho života.

To je umelá inteligencia, o ktorej je táto kniha. Nie povstanie robotov, nie nejaká všeobecná umelá inteligencia zo sci-fi filmu, ale tichý algoritmus v počítači bankovej pracovníčky, ktorý o vás rozhoduje, kým usrkávate kávu zadarmo. Kým sa môžeme baviť o tom, ako Európa takéto systémy reguluje, potrebujeme vy aj ja spoločný, úprimný obraz toho, čím vlastne sú. To je celá úloha tejto kapitoly. Zatiaľ žiadne právo, sľubujem. Len mašinéria.

Dva spôsoby, ako postaviť šachistu

Povedzme, že chceme vy a ja postaviť počítačový program, ktorý hrá šach. Máme dva zásadne odlišné spôsoby, ako na to ísť.

Prvý spôsob je naprogramovať ho explicitne. Sadneme si a napíšeme dlhý zoznam pravidiel: prioritne choď po strelcoch, kontroluj stred a tak ďalej. Alebo zájďme ešte ďalej a zakódujeme každú pozíciu, s ktorou sa program môže stretnúť, čím z neho urobíme jeden obrovský vyhľadávací stroj. Vždy skontroluje, kde je, a prečíta si, čo má zahrať. Môže to fungovať, ale viete si predstaviť tú cenu: mohli by sme roky písať pravidlá a stále nám unikne väčšina toho, čo sa na skutočnej šachovnici deje.

Druhý spôsob je nechať stroj učiť sa. Zozbierame dáta, povedzme 10 000 partií, ktoré v minulosti odohrali slušní hráči, a nakrmíme nimi učiaci sa algoritmus. Hovoríme im **tréninové dáta**. Algoritmus tie partie preštuduje a vytiahne z nich vzory. Nikto mu nepovie „chráň si dámu“; príde na to sám z toho, čo víťazní hráči naozaj robili. Ak sú

dáta dosť veľké a dosť dobré a učiaci algoritmus je slušný, skončíme so strojom, ktorý hrá šach bez toho, aby doň niekto šach kedy naprogramoval.

Tento druhý prístup, stavanie systémov, ktoré odvodzujú svoje správanie z dát a nie z ručne písaných inštrukcií, je spôsob, akým sa stavia väčšina modernej AI, od úverového modelu v banke po chatbota vo vašom telefóne. So šachom sa to dialo už v päťdesiatych rokoch. Tento obraz si držte v hlave, lebo takmer všetko, čoho sa AI Act obáva, z neho vyplýva: ak správanie pochádza z dát, potom ten, kto formuje dáta, formuje správanie. Ako ostrá je táto hrana, uvidíme presne do konca tejto kapitoly.

Slovo, na ktorom zákon naozaj visí: odvodzovanie

Teraz malá, ale dôležitá kalibrácia, kým pôjdeme ďalej. Je lákavé povedať „AI znamená samoučiace sa algoritmy“ a nechať to tak. Sám som to povedal a ako intuícia vám to poslúži dobre. Ale zákon na tom svoju definíciu nevešia a ten rozdiel je dôležitý.

Definícia systému AI v Akte (poriadne ju rozpitváme v kapitole 4) sa otáča okolo jedného slova: **odvodzuje**. Tu je jej jadro.

AKT HOVORÍ · ČLÁNOK 3 BOD 1

„systém AI“ je strojový systém, ktorý je dizajnovaný na prevádzku s rôznymi úrovňami autonómnosti, ktorý môže po nasadení prejavovať adaptabilitu a ktorý pre explicitné alebo implicitné ciele odvodzuje zo vstupov, ktoré dostáva, spôsob generovania výstupov, ako sú predpovede, obsah, odporúčania alebo rozhodnutia, ktoré môžu ovplyvniť fyzické alebo virtuálne prostredie;

Zrozumiteľne: to, čo v očiach zákona robí z niečoho systém AI, je, že si svoje výstupy vypracuje zo svojich vstupov. Odvodzuje, namiesto toho, aby mechanicky vykonával pevný recept, ktorý napísal človek. Všimnite si, čo definícia hovorí o učení: systém „môže po nasadení prejavovať adaptabilitu“. Môže. Systém, ktorý sa raz naučil, bol zmrazený a už sa nikdy neučí, je stále jednoznačne systém AI. Samoučenie je spôsob, akým sa väčšina týchto systémov stavia, ale v samotnej definícii je voliteľné: pojem „samoučiaci“ (v slovenskom znení „samovzdelávacie spôsobilosti“) sa v celom nariadení

objavuje presne raz, v jednom odôvodnení, a to len na vysvetlenie tej voliteľnej adaptability.

POZOR

Lákavá línia uvažovania znie: „Náš model bol natrénovaný pred rokmi a už sa neučí, takže to nie je naozaj AI, takže Akt nie je náš problém.“ Toto uvažovanie zlyháva hneď na prvom slove analýzy. Testom je odvodzovanie, nie priebežné učenie.

Prečo som v tomto taký pedant? Lebo som sledoval inteligentných ľudí, ako kráčajú rovno do presne tejto pasce. Slovo *odvodzovanie* si odložte niekam do bezpečia; kapitola 4 ho úplne rozbalí. Nateraz späť k mašinérii.

Trénovanie, nasadenie, používanie

Nech je systém akýkoľvek, jeho život sa delí na dve veľké fázy a hranica medzi nimi má meno, s ktorým sa v tejto knihe budete stretávať neustále.

Najprv prichádza **fáza trénovania**. Tu predložíme tých 10 000 historických šachových partíí a systém sa z nich učí. Je to obdobie zmeny: vnútorné rozhodovanie modelu ešte formujú prúdiace dáta.

V určitom bode prestaneme. Povieme systému: dosť bolo štúdia, tu je nová partia, teraz ju hraj. Od tej chvíle sme vo **fáze používania**. Systém sa už neučí; robí svoju prácu, znova a znova, na prípadoch, ktoré nikdy nevidel.

Hranica medzi nimi, moment, keď systém presunieme z dielne do sveta, je to, čomu hovoríme **nasadenie**.

MÔJ POHĽAD

Je toto dvojfázové rozdelenie zjednodušenie? Áno, a chcem byť v tom úprimný. Existujú algoritmy online učenia, ktoré sa neprestanú učiť nikdy: preštudujú 10 000 partíí a potom sa učia ďalej z každej novej partie, ktorú hrajú, takže hranica sa rozmazáva. Ale to sú okrajové prípady. Na čítanie tejto knihy, a úprimne aj na čítanie nariadenia, vám dvojfázový obraz poslúži dobre: tréningovanie, potom nasadenie, potom používanie.

Ešte jedna poznámka k slovníku, keď už sme tu. V Akte sa organizácia, ktorá systém AI vo svojej činnosti *používa*, nazýva **nasadzujúci subjekt**: banka prevádzkujúca úverový model je nasadzujúci subjekt, odlišný od **poskytovateľa**, ktorý ho postavil. Tieto roly spresníme neskôr; spomínam ich už teraz len preto, aby vás to slovo neprekvapilo.

Zovšeobecňovanie: študent a skúška

Tu je vlastnosť, ktorá oddeľuje dobrý systém AI od zlého, a najlepšie sa vysvetľuje na študentovi.

Predstavte si niekoho, kto sa pripravuje na skúšku. Možnosť jeden: poriadne naštudovať teóriu, chodiť na prednášky, porozumieť predmetu. Možnosť dva, tá lákavá: zohnať otázky z predchádzajúcich termínov skúšky a namemorovať si odpovede. Ak prídu tie isté otázky, memorujúci preletí. Ak sa otázky čo i len trochu zmenia, je stratený, lebo sa nikdy nenaučil podkladové vzory, len konkrétne príklady.

Systémy AI stoja pred presne tou istou križovatkou. Počas tréningovania systém vždy vidí len **vzorku**: našich 10 000 partíí, bankový spis minulých klientov. Dúfame, že sa naučí vzory, ktoré platia pre celú **populáciu**: každú partiu, ktorá sa kedy odohrá, každého klienta, ktorý kedy vojde. Ak ste niekedy študovali štatistiku, myšlienku vzorky verzus populácie spoznáte; je to ona, v kostýme AI.

Keď sa systém naučí zo svojej vzorky prenositeľné vzory, hovoríme, že **zovšeobecňuje**. Šachový stroj, ktorý zovšeobecňuje, hrá dobre aj v partiách, ktoré sa nepodobajú ničomu v jeho tréningových dátach. Stroj, ktorý si svojich 10 000 partíí iba

namemoroval, hrá skvele, presne dovtedy, kým súper neurobí niečo, čo stroj nikdy nevidel.

Kde sa to pokazí, časť prvá: pretrénovanie

Memorovanie namiesto zovšeobecňovania má technický názov: **pretrénovanie** (overfitting). Študent bifiľujúci staré skúškové otázky bol pretrénovaný na svoje tréningové dáta. Náš hypotetický šachový stroj, ktorý sa príliš opiera o presné historické partie, je pretrénovaný tiež.

Urobme to praktickejším, lebo pri modernej AI to prestáva byť problém kvality a stáva sa problémom právnym. Predpokladajme, že natrénujeme model generujúci obrázky na maľbách slávnych umelcov. Ak sa model pretrénuje, potom vo fáze používania nebude tvoriť nové umenie inšpirované tým, čo študoval; bude reprodukovať skutočné diela, ktoré videl počas trénovania. Náš „kreatívny“ produkt teraz znovu vytvára skutočné maľby a máme na krku problém s autorským právom. Chceli sme model, ktorý vstrebal štýly a vzory a vie vygenerovať niečo skutočne nové. Pretrénovaný model nám dáva drahú kopírku.

Takže pretrénovanie je režim zlyhania číslo jeden: systém sa naučil svoje tréningové dáta príliš doslovne. Režim zlyhania číslo dva je subtilnejší a je to ten, ktorého sa AI Act obáva najviac.

Kde sa to pokazí, časť druhá: zaujatosť

Radi si myslíme, že stroje sú nestranné. Zaujatí sme my ľudia: nosíme si kognitívne skreslenia, predsudky, zlé dni. Algoritmus, ktorý nemá psychológiu, je predsa neutrálny?

Nie. A šachový príklad ukazuje prečo, s takmer trápnu jasnosťou.

Predpokladajme, že tých 10 000 tréningových partií pochádzalo len od štyroch hráčov. Náš stroj sa nenaučí hrať šach všeobecne; naučí sa napodobňovať tých štyroch ľudí. Je zaujatý v prospech ich štýlu. A ak mal jeden z tých štyroch konkrétnu opakujúcu sa slabinu, povedzme zvyk kaziť koncovky, stroj tú slabinu verne zdedí. Nikto tú chybu nenaprogramoval. Prišla cez dáta.

Teraz to zväčšite. Historický záznam našej spoločnosti, text, ktorý sme desaťročia písali a digitalizovali, obsahuje hromadu diskriminácie a predsudkov. Natrénujte na tom zázname model generujúci text a model tie vzory zdedí presne tak, ako šachový stroj zdedil zlú koncovku. Natrénujte úverový model na desaťročiach minulých úverových rozhodnutí a zdedí akúkoľvek nespravodlivosť, ktorú tie rozhodnutia obsahovali. Stroj nie je zaujatý, lebo je zlomyseľný; je zaujatý, lebo je zrkadlo. Preto sa zaujatosť tiahne celým AI Actom ako niť a preto sa k nej budeme stále vracáť.

Späť do banky: ako dáta formujú rozhodnutia a ako sa dajú otráviť

Uzavríme kruh a vráťme sa k vašej hypotéke. Tentoraz si vymeňte strany: už nie ste klient, ste banka a chcete postaviť automatizovaný stroj, ktorý schvaľuje úvery. V skutočnosti by takýto model používal veľa vstupov; zjednodušme si to na dva, aby sme videli mechaniku. Každý klient má **príjem** (nízky, stredný alebo vysoký) a **pomer dlhu** (nízky, stredný alebo vysoký).

Nesadli sme si a nenapísali sme schvaľovacie pravidlá. Verní všetkému vyššie sme model natrénovali na historických dátach: minulí klienti, aký mali príjem a pomer dlhu a či splatili. Z tejto vzorky sa model naučil **rozhodovaciu hranicu**: čiaru cez priestor možných klientov. Na jednej strane čiary sedia profily, ktoré historicky splácali: tie sú schválené. Na druhej strane sedia profily, ktoré boli zamietnuté alebo nesplácali: tie sú odmietnuté. Vysoký príjem? Schválené bez ohľadu na pomer dlhu. Stredný príjem s nízkym dlhom? Schválené. Stredný príjem so stredným alebo vysokým dlhom? Zamietnuté. Nízky príjem? Zamietnuté.

Vy, čakajúci svojich dvadsať sekúnd v pobočke, ste jednoducho bod umiestnený na jednu alebo druhú stranu tej čiary.

Teraz vám ukážem útok, lebo je nádherne poučný. Predpokladajme, že niekto chce našu banku oklamať, niekto, koho profil sedí tesne na nesprávnej strane čiary. Svoj vlastný bod posunúť nemôže. Ale čo keby mohol posunúť *čiaru*?

Čiara vznikla z tréningových dát. Takže útočník naverbuje zo svojho okolia asi tridsať ľudí a pošle ich do našich pobočiek zaregistrovať sa ako bežní klienti. Registrácia vyžaduje len základy, meno a adresu. Ale každý z nich podľa inštrukcií dobrovoľne dodá niečo navyše: „Mimochodom, rád by som uviedol svoj príjem, pomer dlhu a

spomenul, že mi v inej banke poskytnú úver.“ Stredný príjem, stredný pomer dlhu, úver inde poskytnutý a splatený. Všetko vymyslené. A tu je tá nepríjemná časť: banka nemá žiadny osobitný dôvod čokoľvek z toho overovať. Títo ľudia o nič nežiadajú. Sú to len dobrovoľné dáta od priateľských nových klientov, takže idú do databázy tak, ako boli zadané.

Potom príde deň, keď banka svoj model znovu natrénuje na nahromadených dátach. Model, ktorý robí presne to, na čo bol postavený, si všimne zhluk klientov so stredným príjmom a stredným pomerom dlhu, ktorí boli inde dobrými dlžníkmi, a ochotne posunie svoju rozhodovaciu hranicu tak, aby takéto profily prepustil. Naučil sa lož, usilovne. Tomu sa hovorí **otrávenie dát** (data poisoning).

Teraz útočník vojde naozaj. Stredný príjem, stredný pomer dlhu: profil, ktorý by poctivý model zamietol. Otrávená čiara sa posunula a úver je schválený. Žiadny server nebol hacknutý, žiadne heslo ukradnuté. Útočník sa modelu vôbec nedotkol; nakímil ho a to stačilo.

Tento príklad milujem, lebo stláča celú kapitolu do jednej scény. Správanie modelu prišlo z dát, takže ovládať dáta znamenalo ovládať rozhodnutie. Zaujatosť prišla potichu, cez záznamy, ktoré nikoho nenapadlo overiť. A všimnite si, že žiadny jednotlivý krok nevyzeral alarmujúco: tridsať registrácií, pár dobrovoľne uvedených údajov, rutinné pretrénovanie. Preto keď AI Act neskôr pre určité systémy vyžaduje veci ako správu dát a odolnosť voči útokom, nie je to byrokratická fantázia. Je to táto banka, táto čiara, táto červená značka na nesprávnej strane.

Teraz máte kompletnú výbavu na zvyšok knihy: systémy, ktoré svoje výstupy odvodzujú namiesto toho, aby sledovali ručne písané pravidlá, zvyčajne natrénované na dátach a potom nasadené do fázy používania; nádej zovšeobecňovania; a dvojicu režimov zlyhania, pretrénovanie a zaujatosť, s otrávením dát ako živou ukážkou toho, aké krehké celé usporiadanie môže byť. Všetko, čo nariadenie robí, je reakcia na túto mašinériu.

Overte si

1. Nemocnica kúpi diagnostický model, ktorý dodávateľ raz natrénoval a ktorý sa už z nových pacientov neučí. V pojmach životného cyklu z tejto

kapitoly je každodenné používanie modelu v nemocnici: A) Fáza tréovania B) Fáza používania C) Nasadenie D) Online učenie

2. Podľa článku 3 bodu 1 AI Actu je kľúčovou schopnosťou, ktorá robí zo systému „systém AI“, to, že: A) Sa po nasadení priebežne učí a aktualizuje B) Zo vstupov, ktoré dostáva, odvodzuje, ako generovať výstupy C) Bol natrénovaný na veľkom množstve dát D) Prekonáva vo svojej úlohe ľudský výkon

3. Model generujúci obrázky natrénovaný na slávnych maľbách začne používateľom reprodukovat presne tie diela. Toto je klasický prípad: A) Zaujatosti B) Otrávenia dát C) Pretrénovania D) Zovšeobecňovania

4. Náš šachový stroj bol natrénovaný len na partiách štyroch hráčov, z ktorých jeden pravidelne kazil koncovky, a stroj teraz kazí koncovky rovnako. Ponaučenie znie: A) Stroje sú inherentne objektívnejšie než ľudia B) Model dedí chyby a vzory svojich tréningových dát C) Viac tréovania vždy odstráni chyby D) Šach je pre systémy AI príliš zložitý

5. Pri útoku na banku tridsať naverbovaných ľudí nikdy o úver nežiadalo, a predsa model začal schvaľovať profily, ktoré predtým zamietal. Prečo útok fungoval? A) Útočník počas používania hackol kód modelu B) Model nebol správne nasadený C) Do tréningovej množiny sa dostali falošné dobrovoľné dáta a posunuli naučenú rozhodovaciu hranicu D) Zamestnanci banky boli podplatení, aby zmenili schvaľovacie pravidlá

Odpovede: 1. B: systém je za hranicou nasadenia a jednoducho robí svoju prácu na nových prípadoch, čo je fáza používania. 2. B: definícia sa otáča okolo odvodzovania, kým adaptabilita po nasadení je len niečo, čo systém „môže prejavovať“, takže samoučenie je voliteľné. 3. C: model si namemoroval tréningové príklady namiesto toho, aby zovšeobecňoval na nové výstupy. 4. B: nikto tú chybu nenaprogramoval; prišla cez vzorku, z ktorej sa model učil. 5. C: útočník otrávil tréningové dáta neoverenými vymyslenými záznamami a pretrénovaný model poslušne posunul svoju rozhodovaciu hranicu, aby vyhovel lži.

Teraz, keď viete, čo tieto systémy sú a ako potichu sa môžu pokaziť, ďalšia otázka je zjavná: prečo sa práve Európa rozhodla, že to potrebuje zákon?

Prečo Európa reguluje AI

Pred časom som požiadal AI generujúcu obrázky, aby mi nakreslila obrázok spájajúci dve témy: umelú inteligenciu a Európsku úniu. Na prvý pohľad odviedla hviezdnu prácu. Vlajka EÚ, kruh hviezd, niečo, čo vyzerá ako žiariaci neurón, a za tým všetkým nádherný kus európskeho vidieka. Keby som vám to na tri sekundy premietol na slajde, prikývli by ste a šli ďalej.

Teraz sa pozrite druhýkrát, skepticky. Prečo sú tam hviezdy dvakrát? Prečo je počet hviezd nesprávny? Prečo rieka v pozadí nikde nezačína a nikde nekončí? Polia vpredu sú pripravené na žatvu; tie hneď za nimi sú holé, akoby ročné obdobia v polovici obrázka vzdali.

A tu je to, na čom záleží: sú to chyby, ktoré by žiadny ľudský ilustrátor nikdy neurobil. Unavený človek by mohol nakresliť jedenásť hviezd namiesto dvanástich. Žiadny človek nenakreslí vlajku dvakrát a riekou odnikiaľ, všetko to krásne zapracuje a s istotou odovzdá. Tie chyby sú *nové* a sú dosť jemné na to, aby vám unikli, pokiaľ ste ich už nehľadali.

Ten jeden obrázok úprimne nesie väčšinu argumentu tejto kapitoly. Rozbaľme ho na štyri dôvody, prečo si myslím, že AI potrebovala vlastný zákon, a potom sa z toho zákona skúsme spolu vykrútiť, len aby sme videli, čo sa stane.

Štyri dôvody, prečo AI potrebuje vlastný zákon

Po prvé: AI robí nové druhy chýb. Existujúce pravidlá pre výrobky a postupy kvality sú postavené okolo ľudských režimov zlyhania: únava, lajdáctvo, podvod. Systémy AI zlyhávajú inak. Zlyhávajú sebavedome, jemne a vo veľkom, spôsobmi, na ktoré nemáme inštinkt. Spomeňte si na úverový model z predchádzajúcej kapitoly, ktorý sa svoje správanie naučil z dát: otrávnate dáta a otrávite každé rozhodnutie po prúde, pričom nikde nič viditeľne „nie je pokazené“. Rozumná odpoveď je presne to, o čo sa AI Act pokúša: zapísať nové režimy zlyhania a požiadať staviteľov aj audítorov, aby ich cielene hľadali. Nemôžete hľadať chybu, ktorú ste nepomenovali.

Po druhé: AI je technológia, nie aplikácia. Chybný obrázok ma stál jednu vetu promptu a pár sekúnd čakania. V inom prostredí tá istá veta vyprodukuje lacnú, presvedčivú lož. Ale nezastavujte sa pri žiadnom jednotlivom zneužití. Ak chcete odhadnúť budúci dosah AI, porovnajte ju s internetom, vrstvou, na ktorej sa stavia všetko ostatné, a nie s ktoroukoľvek jednou populárnou aplikáciou. Technológie s takýmto dosahom regulujeme: letectvo, farmaceutiká, telekomunikácie. Bolo by zvláštne, keby práve táto ostala výnimkou.

Po tretie: automatizácia mieri na nové ciele. Automatizácia je stará ako pluh a zmierili sme sa s ňou. Čo je skutočne nové, je, že sa teraz automatizuje práca bielych golierov, dokonca aj kreatívna práca. Tento druh práce tomu nikdy predtým nečelil a nepripravil sa naň rekvalifikáciou ani, úprimne, priznaním, že sa to deje. Myslím si, že naša neochota priznať si to je sama osebe rizikom. Regulácia môže aspoň formovať, ako rýchlo a akými prostriedkami tá vlna príde.

Po štvrté: v pravidlách je medzera. Môžete sa rozumne spýtať: nemáme už GDPR? Autorské právo? Základné práva? Máme, a všetky sa AI dotýkajú. GDPR upravuje osobné údaje; autorské právo dáva autorom nástroje, keď ich diela pohltia tréningové množiny. Ale ani jedno z nich nereguluje samotný algoritmus. Vezmite GDPR, zjednodušene: získajte platný súhlas na použitie klientových dát na určité účely, ostaňte v rámci tých účelov a nikto vás veľmi ďalej nespochybňuje. Povedzme, že ste ten súhlas zbierali pred rokmi s jednoduchými marketingovými kampaniami na mysli. Odvtedy vaše schopnosti narástli a teraz, v rámci toho istého súhlasu, viete postaviť profilováciu mašinériu, ktorá skutočne riadi, čo vaši zákazníci kupujú. Zaobchádzanie s dátami je v súlade. Na *systém* sa nikto nikdy nepozrel. To je tá medzera.

Toto zákonodarcovia napísali do úplne prvej vety Aktu, nariadenia (EÚ) 2024/1689, účinného od 1. augusta 2024.

Účelom tohto nariadenia je zlepšiť fungovanie vnútorného trhu a podporiť zavádzanie dôveryhodnej umelej inteligencie (ďalej len „AI“) sústredenej na človeka, a pritom zabezpečiť vysokú úroveň ochrany zdravia, bezpečnosti, základných práv zakotvených v charte vrátane demokracie, právneho štátu a ochrany životného prostredia pred škodlivými účinkami systémov AI v Únii a podporovať inovácie.

Zrozumiteľne: EÚ chce, aby sa AI zavádzala, nie zakazovala. Zavádzala však spôsobom, ktorému ľudia môžu dôverovať, s pomenovanými a riadenými škodami. Držte v hlave obe polovice tej vety. Akt sa bežne opisuje ako nepriateľský voči AI; jeho vlastné vyhlásenie účelu hovorí „podporiť zavádzanie“.

„Fajn, tak jednoducho nebudeme používať AI“

Teraz vám poviem o lákavej únikovej ceste, lebo stretávam ju takmer v každej firme, ktorú školím.

Predstavte si banku z predchádzajúcej kapitoly. Postavila samoučiaci sa úverový model, dopočula sa o otrávení dát, prečítala si pár titulkov o AI Acte a niekto na porade povie: *zabudnime na všetky tie učiace sa algoritmy. Postavme dobrý starý expertný systém založený na pravidlách. Tie fungujú celé veky a žiadny „AI Act“ sa na nás vzťahovať nebude.*

Expertný systém založený na pravidlách je presne to, ako znie. Zhromaždíte svojich najskúsenejších kolegov, vytiahnete z nich ich palcové pravidlá (*ak je klient aktívny, a nekontaktovali sme ho tri mesiace, a pravdepodobne bude chcieť tento produkt, pošli e-mail*) a tie pravidlá zakódujete do jedného stroja. Tieto systémy boli pýchou osemdesiatych a deväťdesiatych rokov a mnohé z nich dodnes potichu bežia v bankách, poisťovniach a nemocniciach. Žiadne učenie z dát, žiadne neurónové siete. Len ľudské poznanie, zapísané.

Takže naša banka postaví svoje úverové rozhodnutia takto: ak je kreditné skóre stredné a príjem nízky, zamietni; ak sú obe stredné, schváľ; a tak ďalej. Čisté, schodovité

rozhodovacie hranice nakreslené ľuďmi. A pretože banka chce byť orientovaná na zákazníka, systém vystaví na svojom webe: zadajte príjem a kreditné skóre a uvidíte, či by ste boli schválení. Priateľská bezplatná služba.

A tiež dvere dokorán.

Útočník začne to verejné rozhranie *sondovať*. Posiela umelé kombinácie príjmu a kreditného skóre, stovky, prechádza vstupný priestor po líniah a zaznamenáva každú odpoveď. Zamietnuté, zamietnuté, schválené... schválené, zamietnuté, schválené. Ak web nemá limit na počet požiadaviek, útočník postupne zmapuje presné rozhodovacie hranice systému. Niekedy sa tomu hovorí krádež modelu a všimnite si: fungovalo to dokonale, hoci nikde naokolo nie je žiadne strojové učenie.

Potom príde fáza dva. Útočník, ktorý presne vie, kde hranice ležia, podáva žiadosti skonštruované tak, aby pristáli tesne vnútri „schválené“: hraničné prípady, ktoré stroj nikdy nemal mávnutím prepustiť, prípady, na ktoré sa mal pozrieť človek. Vlastná transparentnosť banky sa stala plochou útoku.

Tá istá rodina trikov má slávnejšieho bratranca. Samojazdiace autá sa spoliehajú na modely rozpoznávania obrazu, aby čítali dopravné značky. Výskumníci pred rokmi ukázali, že sa dá skonštruovať špeciálny *šum*, vzor pre ľudske oči bezvýznamný, vytlačiť z neho malú nálepku a cez noc ju nalepiť na značku STOP. Ráno autá vidia obmedzenie rýchlosti tam, kde vy a ja vidíme STOP. Model nezdiera naše vnímanie; skúma pixely a pixely boli zvolené tak, aby ho oklamali. (Vývojári odvtedy systémy proti tomuto konkrétnemu kúsku otužili, ale lekcia platí.)

Takže, unikla naša banka AI Actu tým, že sa vyhla strojovému učeniu?

Tu musím byť opatrnejší, než býva folklór, a to oboma smermi. Akt nereguluje „strojové učenie“. Článok 3 bod 1 definuje systém AI podľa toho, čo robí: strojový systém, ktorý pre explicitné alebo implicitné ciele *odvodzuje* zo svojich vstupov, ako generovať výstupy, ako sú predpovede, obsah, odporúčania alebo rozhodnutia. Či systém založený na pravidlách pod túto definíciu spadá, je naozaj otázka od prípadu k prípadu. Niektoré expertné systémy uvažujú nad zakódovanými znalosťami spôsobmi, ktoré sa počítajú ako odvodzovanie, a sú v rozsahu pôsobnosti; iné len vykonávajú pravidlá napísané výlučne ľuďmi a vlastné odôvodnenia Aktu hovoria, že tie pokryté byť nemajú.

POZOR

„Postavili sme to na pravidlách, takže AI Act sa na nás nevzťahuje“ nie je bezpečný záver. Ani opačné tvrdenie, že každý stroj s pravidlami je zachytený. O otázke rozhoduje, či systém *odvodzuje*, ako generovať svoje výstupy, od prípadu k prípadu. Kapitola 4 prechádza presne tým, kadiaľ tá čiara vedie, a je to jedna z prakticky najcennejších čiar v celom nariadení.

Nateraz si vezmite hlbší bod, ktorý považujem za dôležitejší než klasifikačnú hru:

riziká sa nikdy nepýtali, akú techniku ste použili. Sondovací útok, zmanipulované schválenia hraničných prípadov, klienti nesprávne posúdení strojom: to všetko sa stalo „bezpečnému“, staromódnemu stroju s pravidlami.

MÔJ POHĽAD

Ak je vašim inštinktom vyhnúť sa regulácii preznačením technológie, právne môžete alebo nemusíte uspieť, ale režimy zlyhania, ktorých sa regulácia obáva, vás budú sledovať tak či tak. Radšej by som najprv postavil ochranné mechanizmy a až potom sa hádal o definíciách, nie naopak.

Kto tú prácu naozaj robí?

Nariadenie je len text, kým ho ľudia nenesú. Takže kto rieši súlad s AI Actom? Vidím okolo tohto zákona formovať sa štyri persóny, a ak čítate túto knihu, aby ste sa zorientovali, kde sa postaviť, táto časť je pre vás.

Kontrolóri, teda audítori. Každá regulácia plodí audítorov; to nie je nič nové. Nové je, že títo audítori musia byť aspoň čiastočne *technickí*, nie čisto právnici. Porovnajte klasickú kontrolu GDPR s kontrolou AI Actu. Podľa GDPR: funguje odhlásenie z marketingu? Lahké: vytvorte testovacieho zákazníka, prihláste sa, kliknite na odhlásiť, počkajte, pozorujte. Taký audit zvládne ktokoľvek dôsledný. Podľa AI Actu musí mať vysokorizikový systém okrem iného správanie odolné voči zlyhaniu (fail-safe), aby keď odumrie komponent, systém degradoval dôstojne, namiesto toho, aby so sebou stiahol vašich zákazníkov. Ako to zauditovať? Idete *do* systému. Skúmate komponenty, záložné

cesty, kód. Hovoríme o niekom, kto vie čítať Python a rozumie tomu, čo sa deje vnútri modelu. S viacerými svojimi firemnými klientmi už presne takýchto hybridov (napoly dátový vedec, napoly audítor) školíme a očakávam, že „AI Act compliance officer“ sa stane bežným názvom pracovnej pozície.

Stavitelia, teda vývojári a dátoví vedci. Akt ťahá vývojárov do tímu pre súlad tak, ako letectvo kedysi dávno vtiahlo konštruktérov lietadiel: človek, ktorý formuje komponent, musí poznať reguláciu, ktorú ten komponent musí splniť. Dátoví vedci na to nie sú zvyknutí. Doteraz väčšina pracovala, akoby regulácia bola oddelenie niekoho iného. Zvýšiť ich kvalifikáciu tak, aby ich systémy vychádzali v súlade *už z konštrukcie*, a nie záplatované dodatočne, je druhá požiadavka, ktorú od firemných klientov stále dostávam.

Tvorcovia pravidiel a vymáhateľa. Keď som túto látku učil prvýkrát, tieto roly boli väčšinou na papieri. Už nie sú. Úrad pre AI na úrovni EÚ existuje, mašinéria správy a riadenia aj režim sankcií sa uplatňujú od 2. augusta 2025 a každý členský štát musí zriadiť aspoň jedno národné *regulačné experimentálne prostredie* (sandbox): dozorované prostredie, kde môžete v spolupráci s orgánom vyvíjať a testovať inovatívny systém AI skôr, než ho uvediete na trh (článok 57). Termín pre tieto národné experimentálne prostredia posunula novela omnibus z roku 2026 na 2. augusta 2027, takže očakávajte, že táto persóna bude ešte roky rásť.

STAV K AUGUSTU 2026

Stropy pokút sú dramatické, až 35 miliónov EUR alebo 7 percent celosvetového obratu za najhoršie porušenia, a platia od 2. augusta 2025. A predsa, kým toto píšem, nikde neexistuje žiadna potvrdená pokuta podľa AI Actu. Vymáhanie je stroj, ktorý sa ešte skladá, nie kladivo, ktoré už padá.

Strážcovia. AI je hlboko spoločenská téma, takže očakávam, že organizácie občianskej spoločnosti, skupiny pre etickú AI a mimovládne organizácie tu vyrastú na skutočnú silu: budú podávať sťažnosti, testovať systémy zvonku, držať ostatné tri persóny pri poctivosti. Táto je zo štyroch najmenej sformovaná a nebudem predstierať, že predpoviem jej podobu.

Ak chcete moju praktickú radu: prvé dve persóny sú tam, kde sa jednotlivci môžu postaviť už dnes. Buď sa staňte kontrolórom so zmiešanou právno-technickou výbavou, alebo staviteľom, ktorý dodáva systémy v súlade hneď z krabice. Oboch je málo a oboch bude ešte nejaký čas málo.

Optimistický prípad

Túto kapitolu sme strávili útokmi, medzerami a povinnosťami, tak ju uzavriem zámernou zmenou tóniny, lebo budúcnosť s touto reguláciou vnímam skutočne skôr pozitívne.

Odstúpte a spýtajte sa, čo AI Act na najvyššej úrovni vlastne vyžaduje: dôveru, profesionalitu, odolnosť, transparentnosť, zodpovednosť. Ako dátový vedec k vám budem úprimný: náš odbor toto potrebuje. Naše postupy písania kódu, naša testovacia disciplína, naša dokumentácia ešte nie sú tam, kde sedia zrelé inžinierske odbory. Požiadavky Aktu sa do veľkej miery zhodujú s normami vývoja softvéru, ktoré sme si mali osvojiť aj tak. Tímy, ktoré tým prejdú, sa nestávajú len súladnými; stávajú sa lepšími.

Základným kameňom je dôvera. Každé ráno sa spoliehate na svoje auto bez toho, aby ste mu kontrolovali brzdy, lebo celý regulačný ekosystém tú dôveru urobil racionálnou. Naliehanie Aktu na testovanie, riadenie rizík a transparentnosť kladie rovnaké základy pre AI. A odolnosť nie je ani dobročinnosť voči regulátorovi: spýtajte sa sami seba, čo jedno škodlivé zlyhanie (jedna banka, ktorej úverový stroj bol verejne oklamáný) stojí celé odvetvie na erodovanej dôvere.

Je tu aj strategická cena. Predstavte si o dvadsať či tridsať rokov malé zariadenie vo vašej obývačke, ktoré riadi polovicu vášho života: vašu prácu, váš kalendár, rady k vášmu osobnému rozvoju. Chceli by ste mu dôverovať? Teraz si predstavte, že nesie malý štítok, **AI z EÚ**, skratku pre *postavené podľa najprísnejších pravidiel bezpečnosti a etiky AI na svete*. Ten štítok dnes neexistuje a špekulujem s vami. Ale „vyrobené v EÚ“ si presne taký význam vyslúžilo v potravinách a strojárstve a nevidím dôvod, prečo by AI nemohla nasledovať.

Každá veľká technológia prechádza fázou dozrievania. Raný internet bol experimentovanie takmer bez pravidiel; dnešný internet beží na sofistikovaných

protokoloch a bezpečnostnej mašinérii, ktorú dnes nikto nenazve voliteľnou. AI kreslí ten istý kurz, len rýchlejšie. AI Act je stávka Európy, že dozrievanie sa dá riadiť, a nie len prežiť. Niektoré prekážky pred nami sú skutočné a kritiku vám budem cestou ukazovať. Ale samotná stávka je, myslím, rozumná.

Overte si

- 1. AI vygenerovaný obrázok EÚ (dvojité vlajky, nemožná rieka) bol použitý na ilustráciu ktorého argumentu pre reguláciu AI?** A) Systémy AI sú pre malé firmy príliš drahé B) AI robí nové chyby, jemne zapracované, ktoré by žiadny ľudský autor neurobil C) Generátory obrázkov porušujú autorské právo D) Systémy AI na fungovanie vždy vyžadujú osobné údaje
- 2. Prečo samotné GDPR nezatvára regulačnú medzeru okolo AI?** A) GDPR sa vzťahuje len na firmy mimo EÚ B) GDPR bolo zrušené, keď AI Act nadobudol účinnosť C) GDPR upravuje zaobchádzanie s osobnými údajmi, ale neskúma správanie algoritmu; súladný súhlas môže stále kŕmiť škodlivý profilovací systém D) GDPR neobsahuje žiadne sankcie, takže sa nedá vymáhať
- 3. Banka nahradila svoj samoučiaci sa úverový model ručne postaveným expertným systémom založeným na pravidlách a vystavila ho na webe. Čo sa stalo a čo to ukazuje?** A) Nič, lebo systémy založené na pravidlách sa nedajú napadnúť B) Útočník sondoval verejné rozhranie, zmapoval rozhodovacie hranice a zneužil hraničné prípady; riziká nezáviseli od toho, či sa použilo strojové učenie C) Systém bol podľa AI Actu okamžite zakázaný D) Systém prestal fungovať, lebo pravidlá nemôžu robiť úverové rozhodnutia
- 4. Vzťahuje sa AI Act na ten úverový systém založený na pravidlách?** A) Áno, každý softvérový systém používaný v bankovníctve je automaticky systém AI B) Nie, Akt pokrýva len samoučiace sa systémy C) Záleží na tom, či systém *odvodzuje*, ako generovať svoje výstupy; pravidlá, ktoré len vykonávajú ľuďmi napísanú logiku, pokryté byť nemajú, kým expertné systémy, ktoré uvažujú, môžu byť D) Len ak ho banka dobrovoľne zaregistruje
- 5. Ktoré tvrdenie o rolách pri súlade s AI Actom a o vymáhaní je k polovici roka 2026 presné?** A) Audítori budú potrebovať zmiešané právno-technické

zručnosti; režim sankcií (až 35 miliónov EUR alebo 7 percent obratu) sa uplatňuje od augusta 2025, hoci nikde zatiaľ neexistuje potvrdená pokuta B) Systémy AI smú auditovať len právnici a stovky pokút už boli uložené C) Vývojári sú od povinnosti súladu oslobodení; zodpovední sú len manažéri D) Národné experimentálne prostredia sú povinné od roku 2024 a každý členský štát už jedno prevádzkuje

Odpovede. 1 je B: chyby obrázka (zdvojené hviezdy, rieka odnikiaľ) sú režimy zlyhania špecifické pre AI, a preto Akt žiada staviteľov a audítorov, aby hľadali chyby špecifické pre AI. 2 je C: GDPR reguluje zaobchádzanie s dátami a súhlas, nie systém postavený na nich, čo je medzera, ktorú AI Act zaplňa. 3 je B: sondovací útok (krádež modelu) zmapoval cez verejné API ručne nakreslené rozhodovacie hranice, čím dokázal, že nebezpečenstvá prežijú zmenu techniky. 4 je C: článok 3 bod 1 definuje systém AI podľa schopnosti odvodzovať, takže systémy založené na pravidlách nie sú ani automaticky zachytené, ani automaticky vyňaté; kapitola 4 kreslí tú čiaru podrobne. 5 je A: stropy pokút platia od 2. augusta 2025 a k augustu 2026 neexistuje potvrdená pokuta, a povinnosť národných experimentálnych prostredí beží do 2. augusta 2027, kam ju posunula novela omnibus z roku 2026.

Teraz viete, prečo zákon existuje. Ďalej si rozložme celý Akt na stôl a pozrime sa, ako do seba jeho kapitoly, úrovne rizika a dátumy zapadajú v kocke.

AI Act v kocke

Predstavte si pondelkové ráno. Váš generálny riaditeľ vám prepošle novinový článok, „Zákon EÚ o AI: pokuty až 35 miliónov eur“, s jednoriadkovým odkazom: „Týka sa nás to?“

Táto kapitola existuje preto, aby ste vedeli odpovedať do obeda. Nie podrobne (zvyšok knihy sú podrobnosti), ale správne, v správnom poradí, bez tých dvoch či troch nedorozumení, ktoré vidím takmer v každom novinovom článku o AI Acte. Berte to ako pohľad z lietadla, kým pristaneme a začneme chodiť ulicami.

Moji kolegovia radi hovoria, že AI Act je dlhá a zložitá regulácia, a majú pravdu. Čo nechcem, je, aby ste svoje porozumenie stavali tehlu po tehle naprieč mnohými kapitolami a tvar budovy uvideli až na konci. Tvar by ste mali vidieť už teraz. Takže tu je: päť myšlienok, jedna mapa, jedna časová os.

Prvá otázka: je to vôbec AI?

Priznanie. Roky som pracoval ako dátový vedec, väčšinou v bankovníctve. Keby ste ma posadili pred jednoduchý model (povedzme taký, ktorý boduje, ktorí zákazníci majú dostať marketingový e-mail) a spýtali sa: „Je toto umelá inteligencia podľa AI Actu?“, nevedel by som vám dať okamžité áno alebo nie. A takéto veci som staval na živobytie.

To nie je falošná skromnosť. Definícia systému AI v Akte je na okrajoch skutočne rozmazaná a okraje sú presne tam, kde žije väčšina firiem. Všetci sa zhodnú, že ChatGPT, Claude a Gemini sú AI. Ale čo logistická regresia, ktorá pomáha rozhodovať o žiadostiach o úver? Starý expertný systém plný ručne písaných pravidiel? Šachový program, ktorý hrá vyhľadávaním v namemorovaných partiách? Niektoré z nich sú vnútri, niektoré vonku a čiara vedie miestami, kam by ju vaša intuícia nedala.

Takže prvý návyk, ktorý chcem, aby ste si vybudovali, je tento: **nedôverujte svojej intuícii o tom, čo sa počíta ako AI.** Ďalšia kapitola je celá venovaná definícii a sľubujem, že vás prekvapí oboma smermi: veci, ktoré ste považovali za „len štatistiku“,

môžu byť pokryté, a veci, ktoré pôsobia ako AI, môžu vypadnúť. Nateraz nechajte otázku otvorenú.

Model verzus systém

Druhé rozlíšenie znie pedantsky a ukáže sa, že na ňom veľmi záleží.

Model AI je kus technológie: samotná natrénovaná vec, váhy, motor. Sám osebe pre zákazníka nič nerobí. **Systém AI** je to, čo okolo modelu postavíte: aplikácia s používateľským rozhraním, integráciami, dátovými tokmi, účelom. Model je motor; systém je auto.

AI Act reguluje oboje, ale oddelene. Väčšina Aktu je o systémoch, o skutočných produktoch používaných vo svete. Menšia časť (kapitola V, ako uvidíme) je o samotných modeloch na všeobecné účely. Na tom záleží, aj keď nikdy nepredávate produkt. Ak ste povedzme open-source prispievateľ publikujúci veľký model na platforme ako Hugging Face, nikomu nedodávate systém, a stále môžete byť v rozsahu pravidiel pre modely.

Každá AI, ktorá sa dotkne EÚ

Otázka rozsahu pôsobnosti má milosrdne krátku odpoveď. Ak vyvíjate AI vnútri EÚ, ste pokrytí. Ak dovážate AI do EÚ, ste pokrytí. A ak sedíte úplne mimo EÚ, ale výstup vášho systému sa používa v EÚ (zákazníci v Berlíne, zamestnanci v Madride), ste pokrytí tiež. Článok 2 to oblieka do dlhších viet, ale zhrnutie platí: **každá AI, ktorá sa dotkne EÚ, je regulovaná**. Zemepis vášho sídla nie je úniková cesta, čo je presne lekcia, ktorú sa firmy naučili z GDPR.

Poskytovateľ verzus nasadzujúci subjekt

Ešte jedno rozlíšenie, a je to to, od ktorého sa vaša práca na súlade naozaj začne.

Poskytovateľ vyvíja systém AI alebo model a uvádza ho na trh alebo do prevádzky.

Nasadzujúci subjekt používa systém AI v rámci svojej právomoci; nepostavil ho, prevádzkuje ho. (Ak čítate staršie komentáre, občas uvidíte „používateľ“; konečný Akt hovorí nasadzujúci subjekt, a tak bude aj táto kniha. V hovorenej reči, aj v mojom kurze, sa občas skrúti na „nasadzovateľ“.)

Poskytovatelia nesú väčšinu povinností. Nasadzujúce subjekty ich nesú menej, ale nie nula: byť „len“ nasadzujúcim subjektom vás nedostane mimo Aktu.

A teraz praktický zvrat: vaša firma je pravdepodobne oboje. Vezmite nadnárodnú banku, jeden z našich opakujúcich sa príkladov v tejto knihe. Jej model kreditného skóringu, postavený interne vlastným tímom dátových vedcov: banka je poskytovateľ. Hotový nástroj personálneho oddelenia na triedenie životopisov, kúpený od dodávateľa: banka je nasadzujúci subjekt. Marketingový model doladený z dodávateľovho produktu: záleží, a odpoveď sa môže posunúť. Prvý skutočný krok súladu s AI Actom takmer v každej organizácii, ktorej som radil, je presne toto cvičenie: vypíšte svoje prípady použitia AI a pri každom rozhodnite, poskytovateľ, alebo nasadzujúci subjekt?

Štyri pruhy rizika

Správy AI Act zvyčajne prezentujú ako „pyramídu rizika“. Ten obraz nie je nesprávny, ale chcem, aby ste ho videli ako štyri pruhy, do ktorých môže prípad použitia spadnúť, lebo budete ich kontrolovať v tomto poradí.

Pruh jeden: zakázané praktiky. Kapitola II obsahuje krátky zoznam vecí, ktoré s AI jednoducho nesmiete robiť: manipulatívne techniky, ktoré podstatne narúšajú správanie, zneužívanie zraniteľných skupín, určité použitia biometrie a pár ďalších. Všimnite si slovo *praktiky*. Akt nezakazuje prípady použitia ani celé projekty; zakazuje praktiky. Váš náborový nástroj môže byť ako projekt úplne legálny, kým jedna funkcia zahrabaná vnútri (povedzme odvodzovanie emócií uchádzačov z videa) prekročí zakázanú čiaru. Keď auditujete, nepýtate sa „je tento projekt zakázaný?“, ale „predstavuje niečo vnútri tohto projektu zakázanú praktiku?“

Jeden z týchto zákazov si zaslúži svoje presné znenie, lebo tlač ho tak často prekrúti. Sociálne bodovanie sa zvyčajne opisuje ako zákaz *vlád* hodnotiť občanov. Tu je, čo Akt naozaj zakazuje:

uvádzanie na trh, uvádzanie do prevádzky alebo používanie systémov AI na hodnotenie alebo klasifikáciu fyzických osôb alebo skupín osôb počas určitého obdobia na základe ich spoločenského správania alebo známych, odvodených či predpokladaných osobných alebo osobnostných charakteristík, pričom takto získané sociálne skóre vedie [...] ku škodlivému alebo nepriaznivému zaobchádzaniu [...]

Zrozumiteľne: nikto to nesmie robiť. Text nemenuje vôbec žiadneho aktéra; vláda, banka, poisťovňa aj platforma sú zachytené rovnako, ak bodujú ľudí naprieč kontextmi spôsobom, ktorý vedie k neodôvodnenému alebo nesúvisiacemu škodlivému zaobchádzaniu. Skorý návrh to obmedzoval na orgány verejnej moci; konečný Akt nie. Ak vaša firma stavia „skóre dôveryhodnosti zákazníka“ naprieč kontextmi, zákaz je aj váš problém.

Pruh dva: vysokoriziková AI. Ak sa vyhnete zákazom, ďalšia otázka je, či je váš systém vysokorizikový podľa kapitoly III. Dnu vedú dve cesty. Hlavná je príloha III, zoznam citlivých oblastí použitia: nábor, úverová bonita, vzdelávanie, základné služby a tak ďalej. Model úverového skóringu našej banky je učebnicový prípad; hodnotenie úverovej bonity fyzických osôb sedí priamo v prílohe III. Druhá cesta pokrýva AI konajúcu ako bezpečnostný komponent výrobkov, ktoré už reguluje právo EÚ (strojové zariadenia, zdravotnícke pomôcky; to je cesta cez prílohu I). Vysoké riziko je miesto, kam dopadá skutočná váha Aktu: riadenie rizík, kvalita dát, dokumentácia, logovanie, ľudský dohľad, posudzovanie zhody. Podľa mňa je to jadro Aktu a dostáva najdlhšie kapitoly tejto knihy.

POZOR

Pri väčšine systémov z prílohy III nie je posudzovanie zhody nejaká externá inšpekcia. Článok 43 ods. 2 predpisuje postup „na základe vnútornej kontroly“: vlastné posúdenie poskytovateľom, bez účasti notifikovanej osoby. Nezávislé posúdenie treťou stranou je vyhradené najmä pre určité biometrické systémy a cesty cez výrobky z prílohy I.

Či vlastné posúdenie robí režim ľahším, alebo len osamelejším, je férová otázka; vrátime sa k nej.

Pruh tri: povinnosti v oblasti transparentnosti. Kapitola IV (v podstate jeden článok, článok 50) pokrýva systémy, ktoré nie sú vysokorizikové, ale môžu ľudí zavádzať v tom, s čím majú do činenia. Chatboty sa nesmú vydávať za ľudí; obsah vygenerovaný AI musí byť označený; deep fake musia byť označené; ľudom sa musí povedať, keď sa na nich uplatňuje rozpoznávanie emócií alebo biometrická kategorizácia. Malá kapitola, široký dosah: tento pruh sa dotýka takmer každého, kto prevádzkuje chatbota pre zákazníkov.

Pruh štyri: modely AI na všeobecné účely. Kapitola V reguluje veľké základné modely (motory za ChatGPT, Claudom, Midjourney) na úrovni modelu, s dodatočnou vrstvou povinností pre modely klasifikované ako predstavujúce systémové riziko. Ak takéto modely netrénujete, tento pruh sa vás týka väčšinou nepriamo, cez modely, na ktorých staviate.

Jeden prípad použitia môže sedieť vo viacerých pruhoch naraz, a preto kontrolujete všetky štyri.

Mapa: kapitoly I až XIII

Prijatý Akt, nariadenie (EÚ) 2024/1689, účinné od 1. augusta 2024, je usporiadaný do trinástich kapitol. (Staršie návrhy používali „hlavy“ a iné poradie; ak vás blogový príspevok prevádza „hlavou 8“, opisuje text, ktorý už neexistuje.) Tu je mapa, podľa ktorej sa budete naozaj orientovať:

KAPITOLA	ČO TAM ŽIJE
I	Rozsah pôsobnosti, vymedzenie pojmov, gramotnosť v oblasti AI
II	Zakázané praktiky využívajúce AI
III	Vysokorizikové systémy AI: klasifikácia a povinnosti (väčšina Aktu)
IV	Povinnosti v oblasti transparentnosti (článok 50)

KAPITOLA	ČO TAM ŽIJE
V	Modely AI na všeobecné účely
VI	Podpora inovácií, regulačné experimentálne prostredia
VII	Správa a riadenie: úrad pre AI a vnútroštátne orgány
VIII	Databáza EÚ pre vysokorizikové systémy
IX	Monitorovanie po uvedení na trh a dohľad nad trhom
X	Kódexy správania
XI	Delegovanie právomoci
XII	Sankcie
XIII	Záverečné ustanovenia vrátane článku 113, časovej osi

Ak si zapamätáte jeden navigačný fakt, nech je to tento: **GPAI je kapitola V, hneď za transparentnosťou.** Štyri pruhy, ktoré ste sa práve naučili (zákazy, vysoké riziko, transparentnosť, GPAI), sedia v kapitolách II, III, IV a V, v tomto poradí. Zvyšok je mašinéria okolo nich. Kapitoly I až V sú tam, kde strávite väčšinu čitateľského života; VI až XIII vám povedia, kto to všetko vymáha, kde sa vysokorizikové systémy registrujú a čo nesúlad stojí.

Časová os: tri vlny za nami, tri pred nami

Každé nariadenie EÚ nabieha postupne a nábeh AI Actu už mal dejový zvrát, tak vám dám úprimný, aktuálny kalendár, a nie ten z titulkov roku 2024.

Tri vlny sa už cez nás prehнали:

2. februára 2025: začali sa uplatňovať zakázané praktiky spolu s povinnosťou, na ktorú väčšina kalendárov zabudla: článok 4 o gramotnosti v oblasti AI, ktorý žiada poskytovateľov *aj* nasadzujúce subjekty, aby prijali opatrenia na podporu gramotnosti svojich zamestnancov v oblasti AI. Ak vaša organizácia AI vôbec používa, toto sa na vás vzťahuje už vyše roka.

2. augusta 2025: povinnosti pre modely AI na všeobecné účely, štruktúra správy a riadenia a sankcie.

POZOR

Stropy pokút 35 miliónov eur alebo 7 percent celosvetového obratu za zakázané praktiky sú právne živé od augusta 2025, nie 2026. Vystavenie vymáhaniu nečakalo na „hlavný“ termín.

2. augusta 2026: prišli povinnosti v oblasti transparentnosti podľa článku 50: zverejnenie pri chatbotoch, označovanie syntetického obsahu, označovanie deep fake. Tento dátum sa *neposunul*.

A tu je ten dejový zvrat. Pôvodný článok 113 hovoril, že povinnosti pre vysokorizikové systémy sa začnú uplatňovať od augusta 2026. V roku 2026 EÚ prijala novelu digitálny omnibus (nariadenie (EÚ) 2026/1744, odsúhlasené Parlamentom 16. júna a Radou 29. júna, uverejnené v Úradnom vestníku 24. júla a účinné od 27. júla 2026), ktorá dátumy pre vysokorizikové systémy posunula. Sú to pevné kalendárne dátumy, nie podmienené pripravenosťou noriem:

DÁTUM	ČO SA ZAČÍNA UPLATŇOVAŤ
2. feb. 2025	Zakázané praktiky (článok 5) a gramotnosť v oblasti AI (článok 4): živé
2. aug. 2025	Povinnosti pre modely GPAI, správa a riadenie, sankcie až 35 mil. € / 7 %: živé
2. aug. 2026	Povinnosti v oblasti transparentnosti podľa článku 50 (neodložené)
2. dec. 2026	Nový zákaz pridaný omnibusom: AI generované intímne zobrazenia bez súhlasu a materiál zobrazujúci sexuálne zneužívanie detí
2. aug. 2027	Každý členský štát musí mať funkčné národné regulačné experimentálne prostredie

DÁTUM	ČO SA ZAČÍNA UPLATŇOVAŤ
2. dec. 2027	Povinnosti pre vysokorizikové systémy, cesta cez prílohu III (posunutú z 2. aug. 2026 novelou omnibus z roku 2026)
2. aug. 2028	Povinnosti pre vysokorizikové systémy, cesta cez výrobky z prílohy I (posunutú z 2. aug. 2027)

Pre väčšinu čitateľov tejto knihy je najdôležitejší dátum **2. december 2027**: vtedy mašineria pre vysoké riziko zahryzne do prípadov použitia z prílohy III, kde bude klasifikovaná drvivá väčšina vysokorizikovej AI.

Úprimná správa o stave, polovica roka 2026

Keď som túto látku učil prvýkrát, predpovedal som okolo termínu pre vysokorizikové systémy šialenstvo v štýle GDPR: vybookovaní konzultanti, panický nábor, pochybní experti cez noc. Verziu GDPR z roku 2018 som prežil zvnútra banky a vzorec sa zdal zjavný.

Musím hlásiť, že šialenstvo zatiaľ neprišlo.

STAV K AUGUSTU 2026

Nikde v EÚ neexistuje potvrdená pokuta podľa AI Actu, hoci režim sankcií je živý od augusta 2025. V Úradnom vestníku zatiaľ nebola citovaná ani jedna harmonizovaná technická norma, čo tiež znamená, že certifikát ISO/IEC 42001 je fajn doklad o systéme riadenia, ale *nie je* to súlad s AI Actom. Novela digitálny omnibus je už uverejnená (nariadenie (EÚ) 2026/1744, účinné od 27. júla 2026), takže dátumy v tabuľke vyššie sú tie z vestníka. To všetko je druh stavu, ktorý sa mení; overte si ho znovu, kým sa naň spoľahnete.

A samotný omnibus naučil všetkých lekciu, ktorú by som do osnovy z roku 2024 nedal: termíny EÚ, ktoré som svojim klientom opisoval ako pevné od chvíle, keď Akt nadobudol účinnosť, sa môžu legálne posunúť. O jednu novelu neskôr sa z „augusta 2026“ stal „december 2027“.

Znamená to, že sa môžete uvoľniť?

MÔJ POHĽAD

Zákazy a stropy pokút sú živé dnes, takže pruh jeden nie je voliteľný a nikdy nebol. Povinnosti v oblasti transparentnosti sa uplatňujú od 2. augusta 2026. A práca na vysokom riziku (inventarizácia prípadov použitia, roztriedenie poskytovateľ verzus nasadzujúci subjekt, kontrola prílohy III) trvá v každej skutočnej organizácii mnoho mesiacov. Termín sa raz posunul; nestaval by som plán súladu na nádeji, že sa posunie dvakrát. Radšej by som, aspoň nateraz, počítal s horšou možnosťou: december 2027 je skutočný a pokoj roku 2026 je tichá časť krivky, nie dôkaz, že krivka je plochá.

To je pohľad z lietadla. Teraz pristávame a začíname otázkou, od ktorej závisí všetko ostatné: čo presne sa počíta ako AI?

Overte si

- 1. Zákaz sociálneho bodovania podľa AI Actu sa vzťahuje na...** A) Len orgány verejnej moci B) Orgány verejnej moci a firmy konajúce v ich mene C) Ktoréhokoľvek aktéra, verejného alebo súkromného, ktorý takýto systém uvádza na trh alebo používa D) Len firmy nad určitým prahom obratu
- 2. Nemocnica kúpi od dodávateľa AI nástroj na triedenie pacientov a používa ho na svojich pacientoch. Podľa AI Actu je nemocnica...** A) Poskytovateľ B) Nasadzujúci subjekt C) Distribútor D) Nepokrytá, lebo nástroj nepostavila
- 3. Kde v prijatom Akte sedia pravidlá pre modely AI na všeobecné účely?** A) Kapitola V, hneď za kapitolou o transparentnosti B) Kapitola VIII, za správou a riadením C) Hlava 8, blízko konca Aktu D) Príloha III
- 4. Po novele digitálny omnibus z roku 2026, odkedy sa uplatňujú povinnosti pre vysokorizikové systémy z prílohy III (nábor, kreditný**

skóring a podobné prípady použitia)? A) 2. augusta 2026 B) 2. decembra 2026 C) 2. decembra 2027 D) Kedykoľvek budú dostupné harmonizované normy

5. Ktoré tvrdenie o sankciách a vymáhaní je k polovici roka 2026 správne?

A) Pokuty sa nemôžu ukladať, kým sa v decembri 2027 nezačnú uplatňovať pravidlá pre vysokorizikové systémy B) Stropy pokút 35 mil. € / 7 % sa uplatňujú od augusta 2025, ale nikde zatiaľ neexistuje potvrdená pokuta C) Za zakázané praktiky už bolo uložených viacero veľkých pokút D) Sankcie sa vzťahujú len na poskytovateľov, nikdy na nasadzujúce subjekty

Odpovede: 1: C. Článok 5 ods. 1 písm. c) nemenuje žiadneho aktéra; obmedzenie len na vlády existovalo iba v skorom návrhu, takže súkromné sociálne bodovanie je zakázané rovnako. 2: B. Nemocnica používa systém, ktorý nevyvinula, v rámci svojej právomoci, čo je definícia nasadzujúceho subjektu (a nasadzujúce subjekty stále nesú povinnosti). 3: A. Prijaté nariadenie 2024/1689 používa kapitoly a GPAI je kapitola V, hneď za transparentnosťou v kapitole IV. 4: C. Omnibus posunul dátum pre prílohu III z 2. augusta 2026 na 2. decembra 2027 ako pevný kalendárny dátum, nie podmienený normami. 5: B. Kapitola o sankciách sa uplatňuje od 2. augusta 2025, a predsa k augustu 2026 nikde v EÚ neexistuje potvrdená pokuta podľa AI Actu.

Ďalej otvoríme otázku, na ktorú vám táto kapitola povedala neodpovedať intuíciou: čo sa podľa Aktu naozaj počíta ako AI?

Čo sa počíta ako AI

Niekde vo vašej firme je systém, o ktorom nikto neuvažuje ako o umelej inteligencii. Možno je to bodovací stroj, ktorý oddelenie rizík postavilo pred pätnástimi rokmi, tridsať ručne napísaných pravidiel rozhodujúcich, ktoré faktúry sa označia. Možno je to tabuľka, ktorá priemeruje predaje z minulého kvartálu. A niekde inde je systém, ktorý všetci hrdo nazývajú AI, lebo to tak stálo v brožúre dodávateľa.

AI Act sa o brožúry nestará. Stará sa o jednu definíciu, článok 3 bod 1, a tá definícia je bránou do celého nariadenia. Ak systém nie je „systémom AI“ podľa článku 3 bodu 1, môžete Akt zavrieť a ísť domov, aspoň pri tomto systéme. Ak je, pokračujete po ceste: je zakázaný, vysokorizikový, s minimálnym rizikom? Všetko sa začína tu.

Táto kapitola teda odpovedá na klamlivo jednoduchú otázku: čo sa počíta ako AI? Odpoveď postavíme z troch ingrediencií (autonómnosť, adaptabilita a odvodzovanie) a potom ich poskladáme do oficiálnej definície. Sľubujem konkrétne príklady po celý čas, vrátane pár takých, ktoré vás prekvapia oboma smermi: systémy, ktoré by ste nikdy nenazvali AI a sú pokryté, a systémy, o ktorých by ste prisahali, že pokryté sú, a nie sú.

Autonómnosť: kto naozaj rozhoduje?

Povedzme, že vediete banku a chcete stroj, ktorý rozhoduje, kto dostane hypotéku. Máte dva spôsoby, ako ho postaviť.

V prvej verzii rozhoduje AI sama. Klient požiada; volajme ho Robert. Systém si vytiahne jeho históriu transakcií, vzorce výdavkov, príjem, rizikový rating a povie áno alebo nie. Nikto by sa nehádal, že tomuto hovoríme autonómne. Samostatne rozhoduje o niečom dôležitom v živote človeka.

V druhej verzii AI iba odporúča. Ľudský poradca si prečíta odporúčanie, skontroluje podkladové dáta a urobí konečné rozhodnutie. Táto verzia predsa nie je autonómna; stroj v skutočnosti nemôže nič urobiť?

Nebudte takí rýchli. Keď som pracoval ako dátový vedec v bankovníctve, mali sme pre to, čo nasleduje, názov: ľudský faktor. Model je dobrý. Šesť mesiacov poradca kontroluje jeho odporúčania a zakaždým súhlasí. Pomaly, bez toho, aby padlo akékoľvek rozhodnutie, kontrolovanie ustane. Poradca sa naučí stroju implicitne dôverovať a začne jeho odporúčania mávnutím prepúšťať. Systém bol navrhnutý tak, aby nemal žiadnu autonómnosť, a autonómnosť aj tak získal, ničím exotickjším než ľudskou povahou. (Seká to aj opačným smerom: poradca s mesačnou kvótou na kreditné karty môže model prepísať z dôvodov, ktoré nemajú nič spoločné s Robertovou bonitou. Postaviť človeka za AI je vždy zložitejšie, než to vyzerá.)

Odpoveď Aktu na túto hádanku sedí vo fráze „rôzne úrovne autonómnosti“ v definícii, rozbalenej v odôvodnení 12: určitý stupeň nezávislosti činnosti od ľudskej účasti. Slovo *rôzne* robí veľa práce. Aj skromný stupeň nezávislosti (systém vyprodukuje svoje odporúčanie bez toho, aby ho človek spoluvytváral) tento prvok spĺňa. Ukázaním na človeka v slučke z definície uniknete málokedy.

Málokedy, ale nie nikdy. Usmernenia Komisie k vymedzeniu systému AI (C(2025) 924, prijaté vo februári 2025) potvrdzujú, že autonómnosť je skutočná, nevyhnutná podmienka: systém, ktorý vyžaduje plnú, nepretržitú manuálnu účasť, kde prácu v skutočnosti robí človek a stroj je len nástroj v jeho ruke, mimo definície vypadáva. Latka je nízko, ale existuje.

Adaptabilita: päť systémov, jedna otázka

Druhá ingrediencia je adaptabilita: zhruba schopnosť systému meniť sa podľa rôznych podmienok. Aby ste ju precítili, prejdite so mnou päť systémov a pri každom sa pýtajte: adaptuje sa?

Kalkulačka. Plne deterministická. Zadaťte ten istý vzorec, dostanete tú istú odpoveď, navždy. Nie je adaptívna; najľahší verdikt, aký v tejto kapitole dostaneme.

Web, ktorý si v cookie pamätá vašu tmavú tému. Vyberiete si tmavý režim; pri ďalšej návšteve je tma. Systém sa vám v pedantskom zmysle prispôbil. Ale nenaťahujme predstavivosť; nikto toto nariadenie nepísal s cookies na tému na mysli. Pravdepodobne nie je adaptívny.

Expertný systém s ručne napísanými pravidlami. Spomeňte si na šachový stroj zo začiatku knihy, verziu, kde sme explicitne naprogramovali tisíce pozícií a správny ťah pre každú z nich, a hranie znamenalo vyhľadať aktuálnu pozíciu v tej knižnici. Týmto súborom pravidiel sa hovorí expertné pravidlá, lebo ich, dúfajme, napísal expert; prístup mal svoje zlaté časy v osemdesiatych rokoch, okrem iného v medicínskej diagnostike. Je takýto systém adaptívny? V ktoromkoľvek okamihu nie: je zmrazený presne tak, ako ho tvorcovia nechali. Ľudia môžu pravidlá každý týždeň revidovať, ale samotný systém sa používaním nemení.

Systém strojového učenia. To je ten druhý šachový stroj, ktorý sme neprogramovali ťah po ťahu, ale natrénovali na desaťtisíc historických partiách a nechali ho, nech si sám vypracuje, ako hrať. Počas tréningovania sa tento systém adaptuje neustále: skúsi otvorenie, z dát zistí, že existujú lepšie, a zmení stratégiu. Po nasadení sa typický ML systém prestane učiť, ale postavený bol adaptáciou.

Systém strojového učenia s online, priebežným učením. Ten istý šachový stroj, ibaže sa učí aj po nasadení a aktualizuje sa po každej partii, ktorú odohrá. Je to skutočná technika (používa ju predpovedanie počasia, lebo na včerajších dátach záleží), ale okrajová, a z dobrého dôvodu: je krehká. Ak útočník vie, že váš šachový stroj sa učí z každého súpera, môže zámerne hrať hrozné partie a otráviť ho. Možno si pamätáte príbeh o otrávených tréningových dátach zo začiatku knihy; pri online učení sa otrava nikdy nemusí skončiť.

Škála adaptability, od nuly pri kalkulačke po neustálu sebamodifikáciu online učiaceho sa systému. A teraz pointa: pre definíciu Aktu na takmer ničom z toho nezáleží. Definícia hovorí, že systém AI „môže po nasadení prejavovať adaptabilitu“. *Môže.* Budem k vám úprimný: keď som to čítal prvýkrát, strávil som nad tým týždne, písal som o tom články, dokonca som kontaktoval Európsku komisiu, lebo pri prísnom čítaní so všetkými tými „a“ okolo sa zdalo, že žiadny pevný systém sa nemôže kvalifikovať. Ustálená odpoveď, potvrdená usmerneniami z februára 2025, je, že adaptabilita je voliteľná. Ak sa váš systém po nasadení ďalej učí, fajn; ak sa už nikdy nezmení, stále môže byť systémom AI. ChatGPT, ktorý je po tréningu v podstate zmrazený, jeho neadaptívnosť nezachraňuje, a ani váš model.

POZOR

Občas budete počuť AI definovanú ako „samoučiace sa algoritmy“. S touto skratkou opatrne: Akt nedefinuje AI samoučením. Fráza sa v celom nariadení objavuje presne raz, v odôvodnení 12 (v slovenskom znení ako „samovzdelávacie spôsobilosti“), a len na vysvetlenie tohto voliteľného prvku adaptability. Ukotvite sa na samoučení a nesprávne vyradíte z rozsahu každý nasadený model, ktorý sa už neučí, čo je väčšina z nich.

Odvodzovanie: srdce definície

Ak je autonómnosť nízka latka a adaptabilita voliteľná, čo naozaj oddeľuje AI od bežného softvéru? Odôvodnenie 12 vám to hovorí priamo: „Kľúčovou charakteristikou systémov AI je ich spôsobilosť odvodzovať.“ Odvodzovanie je miesto, kde si definícia zaslúži svoju existenciu, tak spomaľme.

Odvodzovanie v zmysle Aktu je spôsobilosť systému derivovať, ako generovať svoje výstupy: vypracovať si modely, vzory alebo pravidlá z dát alebo zakódovaných poznatkov, namiesto toho, aby len vykonával recept, ktorý človek celý vypísal. Náš natrénovaný šachový stroj odvodzuje: nikto mu nepovedal, ako reagovať na pozíciu, ktorú nikdy nevidel, a predsa vyprodukuje rozumný ťah, lebo počas tréningu si vypracoval vlastný model hry. Hypotekárny model odvodzuje: naučil sa z historických klientov, ktorí splatili alebo nesplatili, a naučené vzory aplikuje na Roberta, ktorého nikdy nestretol. ChatGPT odvodzuje: takmer určite nikdy nevidel váš presný prompt, a predsa vygeneruje odpoveď zo vzorov naučených naprieč tréningovými dátami. Odporúčací systém Netflixu alebo Spotify odvodzuje: ste jedineční, ale ľudia podobní vám sú na platforme roky a systém sa od nich naučil, čo by sa vám mohlo páčiť.

A teraz kľúčová hranica, a tu sa neopatrné čítania Aktu mýlia. Odôvodnenie 12 sa nezastaví pri „spôsobilosti odvodzovať“. Pridáva obmedzujúcu vetu, ktorá si zaslúži vlastný rám na stene.

Schopnosť systému AI odvodzovať presahuje základné spracúvanie údajov tým, že umožňuje učenie sa, odôvodňovanie alebo modelovanie.

Zrozumiteľne: nie každý výpočet, ktorý zo vstupu vyprodukuje výstup, je odvodzovanie. Ak štatistik urobí prieskum medzi 500 ľuďmi o ich veku a vypočíta priemer, aby odhadol priemer celého mesta, je to úplne úctyhodný kus inferenčnej štatistiky, a *nie je* to AI podľa Aktu. Usmernenia z februára 2025 to hovoria takmer doslova: softvér, ktorý z prieskumu vypočíta priemer populácie, je ich vlastný príklad základného spracúvania dát, mimo definície. Tá istá veta odôvodnenia vylučuje vašu tabuľku, vaše databázové dopyty, váš mesačný dashboard KPI. Štatistická inferencia a odvodzovanie podľa článku 3 bodu 1 zdieľajú slovo, nie rozsah. Akt ide po systémoch, ktoré sa *učia*, *odôvodňujú alebo modelujú*, po systémoch, ktoré si derivujú vlastný spôsob produkovania výstupov, nie po systémoch, ktoré aplikujú pevný vzorec, nech znie akokoľvek štatisticky.

K odvodzovaniu vedie okrem strojového učenia aj druhá cesta a odôvodnenie 12 ju pomenúva: „prístupy založené na logike a vedomostiach, ktoré odvodzujú riešenie úloh zo zakódovaných poznatkov alebo symbolického znázornenia“. Tú vetu si podržte: o chvíľu rozhodne o osude našich expertných systémov a je selektívnejšia, než sa na prvý pohľad zdá.

Definícia, poskladaná

Teraz máte všetky tri ingrediencie, takže tu je najzávažnejšia veta nariadenia, článok 3 bod 1.

„systém AI“ je strojový systém, ktorý je dizajnovaný na prevádzku s rôznymi úrovňami autonómnosti, ktorý môže po nasadení prejavovať adaptabilitu a ktorý pre explicitné alebo implicitné ciele odvodzuje zo vstupov, ktoré dostáva, spôsob generovania výstupov, ako sú predpovede, obsah, odporúčania alebo rozhodnutia, ktoré môžu ovplyvniť fyzické alebo virtuálne prostredie.

Zrozumiteľne: strojový systém (ľudia robia rozhodnutia pokrytí nie sú; vaša recepcná odporúčajúca turistické atrakcie nie je regulovaný systém), s aspoň nejakou nezávislosťou od ľudskej účasti (nízka latka, ale skutočná), ktorý sa po nasadení môže alebo nemusí ďalej učiť (skutočne voliteľné) a ktorý odvodzuje, ako generovať svoje výstupy (jadro veci, a pamätajte: odvodzovanie musí presahovať základné spracúvanie dát).

Definícia je široká, a zámerne. Technickí ľudia, s ktorými hovorím, hovoria, že ide o odolnosť voči budúcnosti: Komisia nechcela Akt novelizovať zakaždým, keď sa objaví nová architektúra. Moji právni kolegovia pridávajú druhý motív: šírka núti firmy naozaj inventarizovať svoje systémy a nad každým sa zamyslieť, namiesto predpokladu, že nič z toho sa na ne nevzťahuje.

Široká však neznamená bez hraníc. Kadiaľ presne hranica vedie a ktoré systémy podľa Komisie odvtedy padajú mimo nej, je vecou ďalšej časti.

Takže ktoré systémy sú pokryté?

Vynesme verdikt nad našimi piatimi systémami, tentoraz úprimne, vrátane chaotického stredu.

Kalkulačka: vonku. Žiadne odvodzovanie, žiadna debata. **Cookie na tmavú tému:** vonku. Pamätá si preferenciu; neučí sa, neodôvodňuje a nemodeluje nič.

Systém strojového učenia (naš natrénovaný šachový stroj, hypotekárny model): vnútri. Vlastné pravidlá si derivoval z dát; presne pre tento systém bola definícia napísaná. **Online priebežne sa učiaci systém:** dôrazne vnútri. Odškrtať každú kolónku vrátane tej voliteľnej.

A **expertný systém s ručne napísanými pravidlami:** prípad, pri ktorom najviac potrebujem, aby ste to mali správne, lebo európske odvetvia sú ich plné. Bankovníctvo, poisťovníctvo, výroba, všade bežia desaťročia staré stroje s pravidlami a ich vlastníci potrebujú správnu odpoveď, nie slogan. Odpoveď znie: *záleží na tom, či systém odôvodňuje.*

Odôvodnenie 12 kreslí čiaru samo. Na jednej strane výslovne vylučuje „systémy, ktoré sú založené na pravidlách vymedzených výlučne fyzickými osobami na automatické vykonávanie operácií“. Pevný súbor pravidiel, ktorý napísal človek a stroj ho len vykonáva (ak suma presiahne 10 000 a krajina je na zozname, označ faktúru), nie je systém AI, nech je pravidiel koľkokoľvek. Náš šachový stroj s vyhľadávacou tabuľkou patrí sem tiež: každý ťah bol vopred naprogramovaný ľuďmi; stroj vykonáva, neodvodzuje. Usmernenia dokonca používajú šachový program s pravidlami a heuristikami ako vlastný príklad klasickej heuristiky mimo definície.

Na druhej strane odôvodnenie 12 zahŕňa prístupy založené na logike a vedomostiach, ktoré *odvodzujú zo zakódovaných poznatkov*. Skutočný expertný systém v tradícii medicínskej diagnostiky osemdesiatych rokov, teda zakódované expertné poznatky plus odvodzovací stroj, ktorý reťazí pravidlá, dedukuje závery a derivuje odpovede na prípady, ktoré nikto explicitne nenaprogramoval, odvodzuje a je systémom AI. Usmernenia uvádzajú expertné systémy na medicínsku diagnostiku presne ako tento druh prípadu v rozsahu pôsobnosti.

Takže ak inventarizujete dedičstvo starých pravidiel, otázka pri každom systéme znie: len vykonáva ľuďmi napísané pravidlá, alebo odôvodňuje nad zakódovanými poznatkami, aby dospel k záverom, ktoré jeho autori nikdy nevypísali? Vykonávanie je vonku; odôvodňovanie je vnútri.

Čiara pri expertných pravidlách je jedným prípadom širšej korekcie a tu som časom aktualizoval vlastnú radu. Keď bol Akt nový, moje úprimné odporúčanie bolo predpokladať, že všetko s odvodzovaním je v rozsahu, a nehládať východy, lebo žiadne oficiálne zúženie neexistovalo. Od februára 2025 existuje. Usmernenia Komisie k definícii pomenúvajú celé kategórie systémov, ktoré sedia mimo definície, hoci počítajú a v úzkom zmysle dokonca odvodzujú: systémy na zlepšenie matematickej optimalizácie (vrátane dlho zavedených metód ako lineárna a logistická regresia používaných tým ustáleným spôsobom), základné spracúvanie dát, klasické heuristiky a jednoduché predikčné systémy, ktorých výkon by vyrovnalo základné štatistické pravidlo (povedzme predpoveď zajtrajšieho dopytu historickým priemerom). Už nie ste povinní klasifikovať každý vzorec v budove ako AI.

Usmernenia k definícii z februára 2025 (C(2025) 924) ostávajú platným čítaním článku 3 bodu 1 zo strany Komisie. Akt vykladajú, nie novelizujú, a žiadny súd zatiaľ o okrajoch definície nerozhodol, takže sudca by mohol hraničný prípad vidieť inak. Kým sa pri dôležitom systéme spoľahnete na jednu z kategórií mimo rozsahu, overte, či usmernenia neboli aktualizované.

MÔJ POHLAD

Niektoré systémy budú na čiare vykonávanie verzus odôvodňovanie sedieť nepohodlne a je to posúdenie od prípadu k prípadu. Tam, kde systém čiaru naozaj obkročmo sedlá, by som radšej, aspoň nateraz, počítal s horšou možnosťou a zdokumentoval prečo. Zaobchádzať s hraničným systémom ako so systémom AI vás stojí málo; nesprávne ho vyradiť z rozsahu vás môže stáť veľa.

S hranicou v ruke prejdite typickú firmu a stále nájdete v rozsahu dosť: marketingový stroj predpovedajúci, ktorý produkt ponúknuť ktorému zákazníkovi a kedy; modely retailových úverov a poistného; web, ktorý sa preusporadúva podľa naučeného správania návštevníkov; detekciu podvodov pri kartových transakciách; predpovedač výroby solárnej elektriny; systém v aute čítajúci značky STOP; Face ID vo vašom telefóne; prevádzkový model predpovedajúci, ktorý pevný disk zlyhá ako ďalší; a samozrejme každého chatbota od ChatGPT po Claude. Všetko strojové, všetko odvodzujúce.

Jeden hlboký nádych, kým v panike pošlete inventár predstavenstvu: byť systémom AI neznamená byť silno regulovaný. Znamená to, že čítate Akt ďalej. Väčšina systémov z toho zoznamu pristane v teritóriu minimálneho rizika, bez povinností, pre ktoré by sa oplátilo strácať spánok. Niektoré si vyzbierajú povinnosti v oblasti transparentnosti, pár bude vysokorizikových a hárka praktík je priamo zakázaná. To triedenie je to, kam smeruje zvyšok tejto knihy. Definícia je len brána, ale teraz viete, presnejšie než väčšina oddelení pre súlad, kto cez ňu musí prejsť.

Overte si

1. Ktorý prvok definície v článku 3 bode 1 je skutočne voliteľný? A) Byť strojovým systémom B) Autonómnosť C) Adaptabilita po nasadení D) Spôsobilosť odvodzovať

2. Váš tím postaví softvér, ktorý urobí prieskum medzi 500 zamestnancami a vypočíta ich priemerný čas dochádzania, aby odhadol priemer celej firmy. Podľa Aktu a usmernení Komisie k definícii z februára 2025 je to: A) Systém AI, lebo inferencia zo vzorky na populáciu je jadrový koncept Aktu B) Nie systém AI, lebo je to základné spracúvanie dát, ktoré musí spôsobilosť odvodzovať presahovať C) Systém AI, ale vyňatý, lebo sa používa interne D) Zakázaný, lebo spracúva dáta zamestnancov

3. Ktorý z týchto systémov založených na pravidlách je podľa Aktu najpravdepodobnejšie systémom AI? A) Šachový program, ktorý pre každú pozíciu vyhľadá vopred naprogramované ťahy B) Filter faktúr vykonávajúci tridsať pevných ľudmi napísaných pravidiel C) Expertný systém na medicínsku diagnostiku, ktorého odvodzovací stroj reťazí zakódované expertné poznatky, aby dospel k záverom, ktoré jeho autori nikdy explicitne nenaprogramovali D) Termostat, ktorý pod 19 °C zapne kúrenie

4. Hypotekárny model banky vydáva iba odporúčania; konečné rozhodnutie vždy robí ľudský poradca. Pokiaľ ide o prvok autonómnosti v definícii, najbezpečnejšie čítanie je: A) Systém nespĺňa prvok autonómnosti, takže nie je systémom AI B) „Rôzne úrovne autonómnosti“ znamenajú, že sa počíta aj obmedzená nezávislosť, takže človek v slučke systém z definície nevyníma C) Na autonómnosti záleží len pri vysokorizikových systémoch D) Systém je autonómny, len ak poradca prestane kontrolovať jeho výstupy

5. Čo urobili usmernenia Komisie k vymedzeniu systému AI (C(2025) 924, február 2025)? A) Právne novelizovali článok 3 bod 1, aby definíciu zúžili B) Potvrdili, že všetky systémy založené na pravidlách sú systémy AI C) Opísali kategórie systémov mimo definície (ako základné spracúvanie dát, klasické heuristiky a dlho zavedené štatistické metódy) ako nezáväznú čítanie Komisie D) Vytvorili certifikačnú schému na klasifikáciu systémov AI

Odpovede: 1: C, lebo definícia hovorí, že systém „môže po nasadení prejavovať adaptabilitu“, takže systém, ktorý sa po nasadení nikdy nezmení, sa stále kvalifikuje. 2: B, lebo odôvodnenie 12 vyžaduje, aby odvodzovanie presahovalo základné spracúvanie údajov, a usmernenia uvádzajú softvér na priemer z prieskumu ako vlastný príklad mimo rozsahu. 3: C, lebo systémy, ktoré odôvodňujú nad zakódovanými poznatkami, podľa odôvodnenia 12 odvodzujú, kým pevné vykonávanie pravidiel, vyhľadávacie tabuľky a jednoduché heuristiky padajú mimo. 4: B, lebo latka autonómnosti je nízko a na tomto prvku vypadávajú len systémy pod plnou nepretržitou manuálnou kontrolou, takže plánujte, akoby sa odporúčací systém kvalifikoval. 5: C, lebo usmernenia Akt vykladajú, nie novelizujú, pomenúvajú kategórie mimo rozsahu a ako názor Komisie usmerňujú, ale nezaväzujú súdy.

Teraz viete, ktoré systémy prechádzajú prednou bránou Aktu; ďalej stretneme postavy, ktoré im Akt priraduje: poskytovateľov, nasadzujúce subjekty, zamýšľaný účel a zvyšok slovníka Aktu.

Slovník AI Actu

Najprv dobrá správa: najťažšia časť je za vami. V minulej kapitole sme zápasili so samotnou definíciou systému AI: odvodzovanie, autonómnosť, šachový stroj, ktorý sa nakoniec ukázal ako ne-AI. Odteraz sú pojmy kratšie a konkrétnejšie.

Kým sa môžeme baviť o tom, čo Akt zakazuje a čo považuje za vysokorizikové, potrebujeme spoločný slovník. Akt svoje definície nakladá dopredu do článku 3, sú ich desiatky, a hárka nás bude sprevádzať zvyškom knihy ako opakujúce sa postavy. Táto kapitola vyberá tie, ktoré naozaj budete potrebovať. Naučte sa ich teraz a kapitoly o zakázaných praktikách a vysokom riziku sa budú čítať dvakrát rýchlejšie.

Verejne prístupný priestor

Predstavte si dve kamery. Jedna visí v zázemí vašej firmy, kam vkročia iba zamestnanci s kartou. Druhá visí v zákazníckej hale bankovej pobočky, kam môže vojsť ktokoľvek z ulice. Rovnaká kamera, rovnaký softvér, ale pre AI Act žijú v dvoch právne odlišných svetoch, lebo jedna z nich sleduje *verejne prístupný priestor*.

Článok 3 bod 44 ho definuje ako akékoľvek fyzické miesto prístupné neurčenému počtu fyzických osôb. Prácu robia dve veci. Po prvé, „neurčený počet“: vopred neviete, či sa ukáže desať ľudí, alebo tisíc. Po druhé, na vlastníctve nezáleží: súkromne vlastnené nákupné centrum je rovnako verejne prístupné ako mestský park. Odôvodnenie 19 dáva štedrý zoznam: obchody, reštaurácie, kaviarne, banky, posilňovne, štadióny, stanice, letiská, kiná, múzeá, parky, ihriská.

Odôvodnenie 19 kreslí aj výnimky. Kancelárie, výrobné haly a pracoviská určené len pre zamestnancov a poskytovateľov služieb verejne prístupné *nie sú*. Iba odomknuté dvere nerobia z priestoru verejný; ak sú tam značky obmedzujúce prístup, značky vyhrávajú. Zaujímavé je, že online priestory nie sú pokryté vôbec. Váš web alebo fórum nie je „verejne prístupný priestor“, lebo pojem znamená len fyzické miesta. Odôvodnenie končí veľmi európskym zaistením: prístupnosť sa určuje od prípadu k prípadu.

A prečo na tomto pojme záleží? Jeho skutočná váha sedí v jednom konkrétnom zákaze: článok 5 ods. 1 písm. h) zakazuje *dialkovú biometrickú identifikáciu v reálnom čase* vo verejne prístupných priestoroch, ale len *na účely presadzovania práva*, a aj vtedy s úzkymi výnimkami, ktoré stretneme v ďalšej kapitole.

POZOR

Tento pojem sa ľahko prečíta príliš široko. Súkromná banka prevádzkujúca kamerovú analytiku správania zákazníkov vo vlastnej pobočke sa podľa tohto zákazu nedopúšťa zakázanej praktiky. Pokojne môže stavať vysokorizikový systém a GDPR má určite svoje názory, ale zákaz mieri na identifikáciu policajného typu, nie na každú kameru v každej hale.

Pojem si nechajte vo vrecku. Budete ho potrebovať presne raz.

Systém na rozpoznávanie emócií

Predstavte si recepcnú, ktorej zamestnávateľ nainštaloval kameru, ktorá celý deň boduje jej náladu, označuje, keď vyzerá podráždene, a odmeňuje, keď sa usmieva. Ten obraz si podržte; vráti sa v ďalšej kapitole aj s dôsledkami. Nateraz definícia.

Článok 3 bod 39: systém na rozpoznávanie emócií je systém AI na účely identifikácie alebo odvodu *emócií alebo úmyslov* fyzických osôb na základe ich *biometrických údajov*. Biometrické údaje sú zhruba dáta o vašom fyzickom ja: vaša tvár na obrázku, váš hlas, vaša chôdza, vaše oči. Pustite na tieto dáta AI, aby uzavrela „táto osoba je šťastná, nahnevaná, zahanbená“, a postavili ste systém na rozpoznávanie emócií. Odôvodnenie 18 vymenúva emócie, ktoré má na mysli, od šťastia a smútku cez znechutenie a rozpaky po pohrdanie a spokojnosť.

Keď bol oznámený Apple Vision Pro, viedla sa presne o tomto živá diskusia. Senzory milimetre od vášho oka sú vynikajúcim zdrojom biometrických údajov a vaše oko prezradí, čo sa vám páči, aj keď vaša tvár ostane nehybná. Systém odvodzujúci, na ktoré reklamy reagujú vaše zreničky, sedí priamo vnútri tejto definície.

Rovnako dôležité je, čo definícia *vylučuje*: skutočne užitočné systémy môžete stavať až po samú hranicu. Odôvodnenie 18 vyníma dve veci. Po prvé, fyzické stavy ako bolesť alebo únava. Zistiť, že vodič alebo pilot je ospalý, nie je čítanie emócie; je to čítanie fyzického stavu a automobilový priemysel takéto systémy dodáva roky. Po druhé, samotnú detekciu zjavných výrazov, gest alebo pohybov, *pokiaľ* sa nepoužije na odvodenie emócií. Môj obľúbený príklad: kamera v bankovej pobočke s jedinou úlohou, všimnúť si, keď ľudia zdvihnú ruky do vzduchu, slušný skorý signál prebiehajúcej lúpeže. Deteguje gesto a tam sa zastaví. Žiadna odvodená emócia, žiadny systém na rozpoznávanie emócií. Úsmev, zamračené, zvýšený hlas, šepot: detegovať ich ako výrazy je v poriadku. Čiara sa prekročí vo chvíli, keď uzavriete, čo niekto *cíti*.

Prečo na hranici záleží? Lebo odvodzovanie emócií v určitých prostrediach (pracovisko medzi nimi) nie je len vysokorizikové; je zakázané. Naša recepčná sa vráti.

Poskytovateľ a nasadzujúci subjekt

Tu je rozlíšenie, na ktorom visí celý Akt: dve roly, dve knihy pravidiel a tá poskytovateľova je zďaleka ťažšia.

Poskytovateľ (článok 3 bod 3) je každý, kto vyvíja systém AI alebo model AI na všeobecné účely (alebo si ho dáva vyvinúť) a uvádza ho na trh alebo do prevádzky pod svojím vlastným menom alebo ochrannou známkou, či už za odplatu, alebo bezodplatne. Nezáleží na tom, či ste kód napísali vy, alebo ste za to niekomu zaplatili, ani na tom, či si zaň účtujete. Vyvinúť a vydať pod svojím menom, a ste poskytovateľ.

Nasadzujúci subjekt (článok 3 bod 4) je každý, kto systém AI používa v rámci svojej právomoci, s výnimkou čisto osobnej, neprofesionálnej činnosti. Domáci kutil padá mimo Aktu úplne; organizácia používajúca systém AI profesionálne je nasadzujúci subjekt. (Staršie komentáre hovoria „používateľ“. To bolo slovo z éry návrhov; prijaté nariadenie hovorí nasadzujúci subjekt.)

Jedno spresnenie, ktoré nám definície vnucujú: *systém* verus *model*. Vráťte sa k nášmu šachovému stroju. Hotový produkt, rozhranie, webová inštalatérčina a všetko okolo, je systém AI; zahrabaný vnútri sedí komponent, ktorý naozaj hrá šach, model. Model je inteligentné jadro; systém je produkt obalený okolo neho. Rozlíšenie záleží

najviac pri AI na všeobecné účely, kde jedna firma stavia model a stovky stavajú systémy naň.

Dve praktické reality z mojej poradenskej práce. Po prvé, poskytovateľmi je oveľa viac firiem, než očakáva. Je lákavé myslieť si „poskytovatelia znamená OpenAI, Google, Anthropic; my len používame ich veci“. Ale banka, ktorá si natrénuje vlastný model schvaľovania úverov na vlastných dátach klientov, hoci aj v hotovom frameworku, vyvinula systém AI. Je poskytovateľ. Po druhé, jedna firma je zvyčajne *oboje*: veľké organizácie si zmapujú prípady použitia a zistia, že pri piatich sú poskytovateľom a pri troch nasadzujúcim subjektom. Rola sa viaže na prípad použitia, nie na firmu. (Existujú aj dovozcovia, distribútori a splnomocnení zástupcovia; poskytovateľ verzus nasadzujúci subjekt je rozlíšenie, ktoré si treba osvojiť.)

Ked' sa z nasadzujúceho subjektu stane poskytovateľ

A teraz pasca. Kúpili ste od dodávateľa vysokorizikový systém AI; ste nasadzujúci subjekt, s ľahšou knihou pravidiel. Môžete tak ostať? Článok 25 ods. 1 hovorí: len ak od neho dáte ruky preč. Nasadzujúci subjekt, distribútor alebo dovozca sa *stáva* poskytovateľom a dedí plné povinnosti poskytovateľa v troch situáciách:

- umiestnite svoje meno alebo ochrannú známku na vysokorizikový systém, ktorý už je na trhu;
- vykonáte na ňom podstatnú zmenu a ostane vysokorizikový; alebo
- zmeníte zamýšľaný účel systému, ktorý nebol vysokorizikový, spôsobom, ktorý z neho robí vysokorizikový.

Preznačujte ho, prestavajte ho alebo ho prerobte na iný účel do vysokorizikového teritória, a gratulujem: právne je teraz váš; pôvodný poskytovateľ z tej roly vystupuje.

A čo dolad'ovanie? Roky to bola šedá zóna: vezmite predtrénovaný model, dolad'íte ho na vlastných dátach a pýtajte sa, či ste teraz poskytovateľom nového modelu. Od konca roka 2025 existuje pre modely AI na všeobecné účely oficiálna odpoveď. Usmernenia Komisie ku GPAI (C(2025) 7719) hovoria, že úprava existujúceho modelu GPAI z vás robí poskytovateľa GPAI len vtedy, ak je úprava *podstatná*, s orientačným palcovým pravidlom: výpočtový výkon, ktorý miniete, presiahne zhruba tretinu výpočtového

výkonu použitého na tréning pôvodného modelu. To je obrovské množstvo; sama Komisia očakáva, že túto čiaru prekročí len málo dolaďovateľov. Takže typická firma dolaďujúca otvorený model na interných dokumentoch sa nestáva poskytovateľom modelu GPAI, hoci veľmi pravdepodobne stále bude poskytovateľom *systému* AI, ktorý okolo modelu postaví. To sú dve oddelené otázky a článok 25 odpovedá na tú systémovú.

Zamýšľaný účel a odôvodnene predvídateľné nesprávne použitie

Každý systém AI podľa Aktu prichádza s deklarovaným popisom práce. Článok 3 bod 12 ho nazýva *zamýšľaný účel*: použitie, na ktoré systém určil poskytovateľ, vrátane konkrétneho kontextu a podmienok používania, ako sa uvádza v návode na použitie, v propagačných alebo predajných materiáloch a v technickej dokumentácii. Všimnite si tú poslednú časť: váš marketing sa počíta. Ak vaša obchodná prezentácia sľubuje veci, ktoré váš návod vylučuje, rozšírili ste si vlastný zamýšľaný účel.

Konkrétne príklady: „navrhnutý na pomoc lekárom pri diagnostike chorôb pomocou dát pacientov“ znamená, že zamýšľaní operátori sú lekári, nie pacienti. „Určený na vybavovanie zákazníckych dopytov cez chatboty“ znamená chatboty, nie triedenie e-mailov. Každá veta kreslí hranicu.

Teraz predpokladajme, že niekto použije váš systém mimo jeho zamýšľaného účelu a nasleduje škoda. Váš inštinkt: nie môj problém, nasadzujúci subjekt ignoroval návod. Akt tento ťah predvída.

AKT HOVORÍ · ČLÁNOK 3 BOD 13

„odôvodnene predvídateľné nesprávne použitie“ je také použitie systému AI, ktoré nie je v súlade so zamýšľaným účelom, ale ktoré môže byť výsledkom odôvodnene predvídateľného ľudského správania alebo interakcie s inými systémami vrátane iných systémov AI;

Zrozumiteľne: ak ste nesprávne použitie mohli rozumne vidieť prichádzať, očakávalo sa od vás, že s ním budete počítať. Nástroj na pomoc lekárom *skončí* v rukách neškolených ľudí, ktorí sa diagnostikujú sami. To vie predvídať každý, takže poskytovateľ by mal

priať preventívne kroky a vedieť ich preukázať. Chatbot zákazníckeho servisu *bude* provokovaný ku generovaniu nevhodného obsahu, čo je presne dôvod, prečo serióznemu poskytovateľovi prevádzkujú moderačné vrstvy, ktoré dopyty preveria skôr, než model odpovie. „Ale oni to nesprávne použili“ funguje ako obhajoba len vtedy, ak bolo nesprávne použitie skutočne nepredvídateľné. A pri ľudskej kreativite by som zoznam predvídateľných nesprávnych použití držal radšej dlhý než krátky.

Osobitné kategórie osobných údajov

S niektorými dátami je jednoducho nebezpečnejšie počítať. Článok 3 bod 37 definuje *osobitné kategórie osobných údajov* odkazom na GDPR (článok 9 ods. 1) a smernicu o presadzovaní práva: rasový alebo etnický pôvod, politické názory, náboženské alebo filozofické presvedčenie, členstvo v odboroch, genetické údaje, biometrické údaje použité na jedinečnú identifikáciu osoby, údaje o zdraví a údaje týkajúce sa sexuálneho života alebo sexuálnej orientácie.

Jedna nuansa: biometrické údaje sa tu objavujú len v špecifickej príchuti, *na jedinečnú identifikáciu* fyzickej osoby. Odtlačok prsta vás identifikuje jedinečne; to sú údaje osobitnej kategórie. Naša kamera na zdvihnuté ruky pri lúpeži spracúva obrázky tiel, ale nikoho neidentifikuje. Rozhoduje účel, nie pixely.

Palcové pravidlo pre túto knihu: kedykoľvek prípad použitia AI spracúva osobitné kategórie osobných údajov, vaša úroveň rizika stúpa. Niekedy stúpne do vysokého rizika; v kombinácii s určitými praktikami dokonca do zakázaného teritória. Keď auditujete prípad použitia, „dotýka sa tohto zoznamu?“ by mala byť jedna z vašich prvých troch otázok.

Profilovanie

Tu je pojem, na ktorom bude o dve kapitoly ďalej nesmierne záležať; chcem ho zasadiť skoro.

Článok 3 bod 52 definuje *profilovanie* odkazom na článok 4 bod 4 GDPR: akékoľvek automatizované spracúvanie osobných údajov, ktoré hodnotí osobné aspekty fyzickej

osoby, najmä analyzuje alebo predpovedá veci ako jej výkon v práci, majetkové pomery, zdravie, preferencie, spoľahlivosť, správanie, polohu alebo pohyb.

To znie abstraktne, tak to skonkretizujem na našej opakujúcej sa banke.

Predpokladajme, že marketing chce predpovedať, ktorí klienti si čoskoro vezmú hypotéku. Vek medzi 30 a 40, slušný príjem, bývanie v štvrti, z ktorej sa ľudia zvyknú sťahovať: model boduje vaše *majetkové pomery* a životnú fázu. To je profilovanie, učebnicový prípad. Alebo váš zamestnávateľ prevádzkuje softvér, ktorý nielen počíta údery do klávesnice, ale modeluje, akí ste *produktívni*: profilovanie výkonu v práci. Alebo kamera v pobočke boduje pohyby návštevníkov ako „podozrivé“: profilovanie správania.

Všimnite si, aké obyčajné tie príklady sú. Banky, poisťovne a personálne oddelenia to robia desaťročia; Akt tej praxi jednoducho dáva jeden zastrešujúci pojem a pripája dôsledky. Dôsledok, ktorý si treba zapamätať, semienko, ktoré sadím, je tento: keď sa dostaneme ku klasifikácii vysokého rizika, stretnete únikové cesty, ktoré hraničným systémom umožnia vyargumentovať sa z kategórie vysokého rizika. Profilovanie tie cesty zabuchne. Systém z prílohy III, ktorý profiluje fyzické osoby, je *vždy* vysokorizikový, bodka. Keďže profilovanie je taký široký pojem, toto jediné pravidlo zmetie dnu oveľa viac prípadov použitia, než naznačuje intuícia. Pamätajte na hypotekárny model.

Výnimky pre open source

Nakoniec open source. EÚ tu čelila dileme: veľká časť AI ekosystému beží na voľne zdieľaných modeloch a nástrojoch, často udržiavaných ľuďmi, ktorých nikto neplatí, a regulovať ich ako komerčných dodávateľov hrozí uškrténím tej otvorenosti. Akt preto vyrezáva úľavu na dvoch úrovniach a oplatí sa držať ich oddelene.

Úroveň modelu. Poskytovatelia modelov AI na všeobecné účely nesú podľa článku 53 ods. 1 dokumentačné povinnosti: technickú dokumentáciu pre úrad pre AI a vnútroštátne orgány a informačné balíky pre nadväzujúcich staviteľov systémov. Článok 53 ods. 2 presne tieto povinnosti pre otvorené modely sníma.

Povinnosti stanovené v odseku 1 písm. a) a b) sa nevzťahujú na poskytovateľov modelov AI, ktoré sa vydávajú na základe bezplatnej licencie s otvoreným zdrojovým kódom, ktorá umožňuje prístup k modelu, jeho používanie, úpravu a distribúciu, a ktorých parametre vrátane váh, informácií o architektúre modelu a informácií o používaní modelu sú verejne dostupné. Táto výnimka sa nevzťahuje na modely AI na všeobecné účely so systémovými rizikami.

Zrozumiteľne: otvorené vydanie *je* vstupné. Publikujte pod skutočne slobodnou licenciou (takou, ktorá umožňuje prístup, používanie, úpravu a distribúciu), s verejnými váhami, informáciami o architektúre a o používaní, a formálne dokumentačné povinnosti odpadnú. Férové, keďže všetko už je vonku.

Dve povinnosti prežívajú aj pre open-source poskytovateľov a to ľudí zaskočí. Článok 53 ods. 1 písm. c) a d), politika autorského práva a *verejne dostupné zhrnutie obsahu použitého na tréningovanie*, sa vzťahujú na všetkých poskytovateľov GPAI, otvorených aj uzavretých. Ide o zhrnutie (existuje povinná šablóna úradu pre AI), nie o požiadavku odovzdať samotné tréningové dáta. Open-source poskytovateľ preskočí dokumentačné papierovanie, ale zhrnutie tréningového obsahu stále publikuje, ako všetci ostatní.

A výnimka má vlastnú výnimku, v tej poslednej citovanej vete: modely so *systémovým rizikom* nedostávajú žiadnu úľavu. Systémové riziko sa viaže na najschopnejšie modely na špici. Akt ho predpokladá nad 10^{25} operácií s pohyblivou rádovou čiarkou tréningového výpočtového výkonu a poskytovateľa, ktorí túto čiaru prekročia, to musia do dvoch týždňov oznámiť úradu pre AI. Tu záleží na princípe: otvorenosť nevykúpi najväčšie modely z najťažších povinností.

Neexistuje verejný register modelov GPAI označených ako nesúce systémové riziko, takže odolám pokušeniu menovať, a s akýmkoľvek zoznamom, ktorý uvidíte online, by som zaobchádzal podozrievavo. Kým sa na niektorý spoľahnete, overte aktuálny stav v publikáciách úradu pre AI.

Úroveň systému. Oddelene, a ľahko prehliadnuteľné, článok 2 ods. 12 hovorí, že nariadenie sa *vôbec nevzťahuje* na systémy AI vydané na základe bezplatných licencií s otvoreným zdrojovým kódom, pokiaľ nie sú uvedené na trh alebo do prevádzky ako vysokorizikové systémy alebo nespádajú pod článok 5 alebo článok 50. Takže open-source úľava existuje aj pre systémy, nielen modely. Ale trvá len dovtedy, kým sa systém nezatúla do vysokorizikového, zakázaného alebo transparentnostne relevantného teritória. Vydajte open-source nástroj na triedenie životopisov ako produkt a je vysokorizikový ako každý iný; licencia nie je štít.

MÔJ POHĽAD

Konkrétna stratégia vydávania je otázka pre vášho právneho poradcu, takže túto časť berte s obvyklou rezervou. Ale tvar je priateľský: Akt sa skutočne snaží open source nezabiť a zároveň odmieta nechať „otvorené“ stať sa dierou pre najrizikovejšie použitia.

Overte si

1. Súkromne vlastnené kino nainštaluje do svojej haly systém AI. Pre AI Act je hala: A) Nie verejne prístupný priestor, lebo kino je v súkromnom vlastníctve B)

Verejne prístupný priestor, lebo môže vojsť neurčený počet ľudí C) Nie verejne prístupný priestor, lebo sa vyžadujú lístky D) Verejne prístupný priestor len počas otváracích hodín

2. Ktorý z týchto je systémom na rozpoznávanie emócií podľa článku 3 bodu 39? A) Kamera detegujúca, že vodič je unavený B) Kamera detegujúca zdvihnuté ruky v bankovej pobočke na označenie možnej lúpeže C) Systém analyzujúci dáta zo sledovania očí na odvolenie, ktoré reklamy sa osobe páčia D) Mikrofón detegujúci, že volajúci zvýšil hlas, bez vyvodzovania záverov

3. Logistická firma kúpi od dodávateľa systém AI na všeobecné účely, preznačuje ho pod vlastnú ochrannú známku a predáva ďalej ako vysokorizikový nástroj na správu flotily. Podľa článku 25 ods. 1 je firma teraz: A) Stále nasadzujúci subjekt, lebo model netrénovala B) Distribútor bez ďalších

povinností C) Poskytovateľ systému, s plnými povinnosťami poskytovateľa D) Vyňatá, lebo pôvodný dodávateľ si navždy ponecháva všetky povinnosti

4. Open-source model GPAI (bez systémového rizika) je vydaný pod slobodnou licenciou s verejnými váhami, architektúrou a informáciami o používaní. Ktorá povinnosť sa na jeho poskytovateľa stále vzťahuje? A)

Žiadna, lebo open-source modely sú úplne mimo Aktu B) Technická dokumentácia pre úrad pre AI podľa článku 53 ods. 1 písm. a) C) Informačný balík pre nadväzujúcich poskytovateľov podľa článku 53 ods. 1 písm. b) D) Verejne dostupné zhrnutie obsahu použitého na tréningovanie podľa článku 53 ods. 1 písm. d)

5. Model banky predpovedá, ktorí klienti si pravdepodobne vezmú hypotéku, na základe veku, príjmu a štvrte. V slovníku Aktu je to: A)

Profilovanie, lebo automatizovaným spracúvaním osobných údajov hodnotí majetkové pomery fyzickej osoby B) Nie profilovanie, lebo marketing je oprávnený záujem C) Rozpoznávanie emócií, lebo odvodzuje úmysly klientov D) Profilovanie len vtedy, ak klienti výslovne súhlasili

Odpovede: 1: B. Na vlastníctve nezáleží; počíta sa prístup neurčeného počtu ľudí a podmienky ako kúpa lístka to nemenia. 2: C. Únava je fyzický stav a zdvihnutá ruka či hlas sú iba detekcia gesta alebo výrazu, ale odvodzovanie záľub a úmyslov z dát o očiach je presne to, čo definícia pokrýva. 3: C. Umiestnenie vlastného mena alebo ochrannej známky na vysokorizikový systém, ktorý už je na trhu, z vás podľa článku 25 ods. 1 písm. a) robí poskytovateľa. 4: D. Článok 53 ods. 2 vyníma open-source modely bez systémového rizika z písmen a) a b), ale zhrnutie tréningového obsahu a politika autorského práva sa vzťahujú na všetkých poskytovateľov GPAI. 5: A. Automatizované hodnotenie osobných aspektov ako majetkové pomery je profilovanie bez ohľadu na účel.

Teraz, keď zdieľame slovník Aktu, môžeme otvoriť jeho najtemnejšiu zásuvku: praktiky, o ktorých Európa rozhodla, že ich nesmie robiť nikto.

Zakázané praktiky

Pozrite sa na svoj telefón. Je na ňom hra, s ktorou neviete celkom prestať? Taká, ktorá vám presne vo chvíli, keď sa rozhodnete skončiť, podá bonus, šťastnú sériu, postrčenie „ešte jeden level“, ktoré pristane presne v správnom momente? Za tým krásnym frontendom môže byť systém, ktorý sleduje, ako hráte, a individualizuje zážitok, rozhoduje, že *tento* hráč práve teraz potrebuje *tento* boost, aby ďalej míňal. Ten trh som skúmal a môžem vám povedať: takéto systémy existujú.

Zhruba tam sa EÚ rozhodla nakresliť červenú čiaru. Väčšina AI Actu je o riadení rizika: dokumentácia, dohľad, transparentnosť. Kapitola II je iná. Vymenúva praktiky, ktoré sú v Európskej únii jednoducho zakázané, žiadne papierovanie sa neprijíma. A nie je to problém budúcnosti: zákazy sa uplatňujú od 2. februára 2025, nesú najvyššiu úroveň pokút Aktu, až 35 miliónov eur alebo 7 percent celosvetového obratu, a (k tomu sa vrátíme) k augustu 2026 nikto nikde v skutočnosti pokutu nedostal.

Článok 5 ods. 1 obsahuje osem zákazov a deviaty je na ceste. Osem právnych odsekov je suché čítanie, tak ich zoskupujem do štyroch tém: manipulácia a klamanie, zneužívanie zraniteľností, diskriminačné bodovanie a biometrické praktiky. Prejdime si ich tak, ako by som to urobil s klientom: čo je zakázané, čo zakázané len *vyzerá* a kde sa hrany naozaj rozmazávajú.

Manipulácia a klamanie

V klasických detektívkach je zloduch podvodník so šarmom a starostlivo utkanými klamstvami. Dnešná verzia šarm nepotrebuje. Potrebuje vaše behaviorálne dáta a metriku na optimalizáciu. Princíp za touto prvou témou: AI nesmie byť navrhnutá spôsobom, ktorý ľudí oklame alebo tajne ovplyvní k rozhodnutiam škodlivým pre nich alebo pre iných.

Článok 5 ods. 1 písm. a) mieri na AI, ktorá využíva podprahové techniky mimo vedomia osoby alebo manipulatívne či klamlivé techniky, aby narušila správanie niekoho spôsobom, ktorý spôsobuje alebo pravdepodobne spôsobí značnú ujmu. A „značná ujma“ je široká: ekonomická strata sa počíta, psychická ujma sa počíta.

Jedna fráza v Akte si zaslúži vašu plnú pozornosť:

AKT HOVORÍ · ČLÁNOK 5 ODS. 1 PÍSM. A)

„...s cieľom alebo účinkom podstatne narušiť správanie osoby...”

Zrozumiteľne: úmysel sa nevyžaduje. Môžete si v Akte prečítať slovo „účelovo“ a myslieť si, že zákaz zachytáva len vývojárov, ktorí *vedia*, že manipulujú ľudí. Usmernenia Komisie k zakázaným praktikám (február 2025) tieto dvere výslovne zatvárajú: „účelovo“ sa číta objektívne, podľa toho, na čo je technika navrhnutá alebo na čo objektívne mieri. Nikto nemusí ujmu zamýšľať; usmernenia dokonca hovoria, že zákaz pokrýva systémy AI, ktoré manipulujú ľudí *bez toho, aby to ktorýkoľvek človek zamýšľal*, napríklad manipuláciu, ktorú sa systém naučil sám z tréningových dát, lebo náhodou zvyšovala zapojenie. Ak váš odporúčací systém sám objavil, že hnev udrží ľudí pri scrollovaní, „to sme nikdy nemysleli“ nie je obhajoba.

Takže moja rada firemným klientom je jednoduchá: prejdite svoje nasadené prípady použitia a položte dve otázky. Individualizuje tu niečo zážitok spôsobom, ktorý ľudí postrkuje k rozhodnutiam? A ak zákazník postrčenie nasleduje, mohlo by to spôsobiť značnú ujmu? Predstavte si banku, ktorej kampaňový systém sa naučí, že hypotekárna ponuka formulovaná *presne takto*, poslaná *presne* v správnej chvíli, maximalizuje šancu, že finančne neskúsený klient zahryzne. To už nie je marketing. To je algoritmus za marketingom a algoritmus je presne to, na čo tento článok mieri.

Zneužívanie zraniteľností

Druhá téma je ostrejším súrodencom prvej. Podľa článku 5 ods. 1 písm. b) AI nesmie zneužívať zraniteľnosti osoby alebo skupiny z dôvodu ich veku, zdravotného postihnutia alebo osobitnej sociálnej alebo ekonomickej situácie, opäť s cieľom alebo účinkom narušiť správanie smerom k značnej ujme.

Deti sú zjavný prípad. Dieťa hrajúce online hru si ešte len rozvíja kontrolu impulzov; AI, ktorá sa z herných vzorcov naučí psychológiu dieťaťa a potom neúnavne tlačí nákupy v aplikácii, je presne to, čo tento článok opisuje. Ale všimnite si tretiu kategóriu: osobitná sociálna alebo ekonomická situácia. Predstavte si niekoho po zdravotnej núdzi

alebo strate práce, kto zúfalo hľadá peniaze a stretne úverového chatbota, ktorý nesúcíť, ale počíta, cieľi na ľudí v núdzi, lebo štatisticky pravdepodobnejšie prijmú premrštené sadzby. Pohľad Aktu je, že toto nie je zákaznícka interakcia medzi rovnými. Je to mocenská nerovnováha, premenená na zbraň.

Tu je tá zákerná časť: vývojári to často vôbec nenavrhli. Systém optimalizujúci konverziu si zraniteľný segment nájde sám, lebo zraniteľní konvertujú. Preto aj tu „nezamýšľali sme to“ zlyháva: *účinnok* stačí.

Diskriminačné bodovanie

Tretia téma je o AI, ktorá človeka zredukuje na vypočítané skóre a nechá skóre diktovať jeho príležitosti. Akt na to útočí z dvoch uhlov: bodovanie vášho spoločenského života a bodovanie vašej kriminálnej budúcnosti.

Sociálne bodovanie

Ak ste videli epizódu Black Mirror *Nosedive*, tomuto zákazu už rozumiete. Spoločnosť, kde každý nosí neviditeľné skóre poskladané zo spoločenského správania a skóre diktuje, či dostanete byt, prácu, miesto v lietadle.

Článok 5 ods. 1 písm. c) zakazuje AI, ktorá hodnotí alebo klasifikuje ľudí počas určitého obdobia na základe ich spoločenského správania alebo známych, odvodených či predpokladaných osobných charakteristík, keď výsledné skóre vedie k jednej z dvoch vecí: k škodlivému zaobchádzaniu v kontextoch *nesúvisiacich* s tým, kde boli dáta zozbierané, alebo k zaobchádzaniu, ktoré je *neodôvodnené alebo neprimerané* samotnému správaniu. Vaše dovolenkové fotky, ktoré vás stoja pracovný pohovor, sú prvá vetva. Jeden drsný online komentár spred rokov, ktorý vám stále nafukuje poisťné, je druhá.

Dve veci, ktoré ľudia pri tomto zákaze bežne chápu zle. Po prvé, zákaz sa vzťahuje na **ktoréhokoľvek aktéra**, nielen na vlády. Štátny systém čínskeho typu je slávny obraz, ale súkromná poisťovňa stavajúca „skóre zodpovednosti“ z vašich objednávok jedla a hudobného vkusu je priamo v rozsahu tiež. Po druhé, *nezakazuje* bodovanie ako také. Moji typickí klienti sú banky a poisťovne a tie počítajú skóre úverového rizika desaťročia. Model bonity používajúci finančné dáta, vo finančnom kontexte, na

finančné rozhodnutie nie je sociálne bodovanie; je to vysokoriziková AI (počkajte na ďalšiu kapitolu). Problém sa začína, keď vstupy zdriftujú do spoločenského správania a výstupy presiaknu do nesúvisiacich kontextov. Ak váš rizikový model potichu trestá impulzívny nákup spred rokov alebo zhlukuje klientov podľa vzorcov, ktoré sa ukážu ako zástupné znaky toho, s kým sa stýkajú, driftujete k čiare.

Prediktívna policajná práca

Druhá strana tejto témy boduje niečo ešte citlivejšie: či spáchate trestný čin. Poznáte film *Minority Report*, zatýkanie ľudí za zločiny, ktoré ešte nespáchali. Článok 5 ods. 1 písm. d) zakazuje AI, ktorá posudzuje alebo predpovedá riziko, že osoba spácha trestný čin, ale čítajte pozorne kvalifikátor: **výlučne na základe profilovania alebo posúdenia jej osobnostných vlastností a charakteristík**.

To slovo „výlučne“ je nosné, rovnako ako výnimka, ktorá po ňom nasleduje: zákaz sa nevzťahuje na AI, ktorá podporuje ľudské posúdenie, *ktoré je už založené na objektívnych a overiteľných skutočnostiach priamo spojených s trestnou činnosťou*. Takže AI, ktorá vás označí za budúceho zločince kvôli vášmu PSC, osobnostnému profilu a známym: zakázaná. AI pomáhajúca vyšetrovateľom usporiadať dôkazy v existujúcom prípade postavenom na skutočných faktoch: nezakázaná.

POZOR

„Prediktívna policajná práca je úplne zakázaná“ je nesprávne čítanie, ktoré počúvam často, a nechcel by som, aby ste ho opakovali kolegom. Zákaz je užší, než naznačuje filmový plagát: zachytáva predpovede založené *výlučne* na profilovaní alebo osobnostných vlastnostiach, kým AI podporujúca ľudské posúdenie založené na faktoch je výslovne vyňatá.

Biometrické praktiky

Posledná téma je tá, ktorá konzistentne spúšťa poplach, od sci-fi po znepokojujúce titulky: vaša tvár, vaše telo a vaše správanie na verejnosti sa stávajú neustále sledovateľným digitálnym podpisom. Vo verejných priestoroch sa vždy očakávalo isté

splynutie s davom; hromadná biometrická analýza to rozpúšťa. Akt tu zhromažďuje štyri zákazy a sú prísne.

Necielené zbieranie tvárí

Tento je krátky a nezvyčajne konkrétny. Podľa článku 5 ods. 1 písm. e) nesmiete vytvárať ani rozširovať databázy na rozpoznávanie tváre necielenou extrakciou podôb tváre z internetu alebo zo záznamov CCTV. Ak vám firma vysávajúca miliardy fotiek tvárí zo sociálnych sietí, aby predávala služby vyhľadávania podľa tváre, znie povedome, áno: ten biznis model je cieľom. (Slávne pokuty presne pre takúto firmu boli mimochodom pokuty podľa GDPR, nie podľa AI Actu. Držte tie dva režimy oddelene.)

Odvodzovanie emócií v práci a v škole

Predstavte si recepcnú, ktorá trávi celý deň pod náladovou kamerou: softvér boduje jej úsmev, označuje, keď pôsobí podráždene, hlási „zapojenie“ jej manažérovi. Podľa článku 5 ods. 1 písm. f) je ten systém zakázaný. AI sa nesmie používať na odvodzovanie emócií osoby v oblasti pracoviska a vzdelávacích inštitúcií.

Existuje jedna výnimka: systémy zavedené zo zdravotných alebo bezpečnostných dôvodov. Detekcia únavy pre vodičov diaľkovej dopravy alebo operátorov strojov môže stáť na bezpečnostnej nohe. Ale HR dashboard nálady zamestnancov, nástroj v call centre bodujúci nálady operátorov, kamera v triede meriaca pozornosť študentov: to nie sú hraničné prípady. Sú to ťažisko zákazu a podľa mojej skúsenosti sú to aj zákazy, ktorý firemné audity najčastejšie prehliadnu, lebo dodávatelia emočnej analytiky ich predávajú ako neškodné nástroje produktivity.

Všimnite si hranicu: odvodzovanie emócií *mimo* práce a vzdelávania (povedzme na zákazníkoch) zakázané nie je, ale pristane namiesto toho vo vysokorizikovom teritóriu. Zakázané tu ani zďaleka neznamená „všade inde v poriadku“.

Biometrická kategorizácia citlivých znakov

Podľa článku 5 ods. 1 písm. g) AI nesmie kategorizovať jednotlivcov na základe ich biometrických údajov s cieľom odvodiť alebo vyvodiť ich rasu, politické názory, členstvo v odboroch, náboženské alebo filozofické presvedčenie, sexuálny život alebo sexuálnu

orientáciu. Zrozumiteľne: žiadne „táto tvár vyzerá gay“, žiadne „táto chôdza naznačuje protestujúceho“, žiadne odvodzovanie náboženstva z výzoru. Zákaz má technické výnimky (označovanie alebo filtrovanie zákonne získaných súborov biometrických údajov a určité kontexty presadzovania práva), ale jadrová myšlienka je, že vaše telo nie je legitímny vstup na hľadanie najchránennejších faktov o vás.

Diaľková biometrická identifikácia v reálnom čase

A teraz zákaz, o ktorom počul každý, zvyčajne v skomolenej podobe: rozpoznávanie tvárí vo verejných priestoroch.

Tu je, čo článok 5 ods. 1 písm. h) naozaj hovorí. Zákaz pokrýva systémy diaľkovej biometrickej identifikácie **v reálnom čase**, vo **verejne prístupných priestoroch**, používané **na účely presadzovania práva**. Na všetkých troch prvkoch záleží. V reálnom čase znamená živé skenovanie davov, nie detektíva, ktorý neskôr porovnáva fotku. Verejne prístupné priestory znamenajú ulice, parky, dopravné uzly. A „na účely presadzovania práva“ znamená, že tento zákaz je o polícii, nie o firmách.

Aj pre políciu existujú tri kategórie výnimiek: cielené pátranie po obetiach únosov, obchodovania s ľuďmi alebo sexuálneho vykorisťovania a po nezvestných osobách; predchádzanie konkrétnemu, závažnému a bezprostrednému ohrozeniu života alebo skutočnému, predvídateľnému teroristickému útoku; a lokalizácia podozrivých zo závažných trestných činov (činov uvedených v prílohe II s trestnou sadzbou najmenej štyri roky). Aj vtedy každé použitie potrebuje predchádzajúce povolenie súdneho alebo nezávislého správneho orgánu, posúdenie vplyvu na základné práva a registráciu v databáze EÚ. Toto nie je diera, cez ktorú prevezie sledovací štát; sú to úzke dvere s tromi zámkami a každý členský štát vo vnútroštátnom práve rozhodne, či ich vôbec otvorí.

A teraz korekcia, ktorú robím najčastejšie. Nákupné centrum, ktoré páruje tváre návštevníkov s profilmi na sociálnych sieťach, aby stávalo predajné spotrebiteľské spisy, pôsobí, akoby to tu muselo byť pokryté. Nie je.

POZOR

Usmernenia Komisie výslovne hovoria, že použitie rozpoznávania tváří súkromnými aktérmi, v reálnom čase alebo nie, padá mimo tohto zákazu, lebo zákaz je vymedzený na presadzovanie práva. To nerobí scenár s nákupným centrom legálnym alebo bezpečným: čelne sa zráža s pravidlami GDPR o biometrických údajoch a systémy biometrickej identifikácie sú podľa prílohy III vysokoriziková AI. Nesprávna klieťka, tá istá zoo. Ale ak proti nákupnému centru budete citovať článok 5, jeho právnici tú výmenu vyhrajú.

Deviaty zákaz: AI generované intímne zobrazenia a materiál zobrazujúci zneužívanie detí

Zoznam sa rozšíril. Novela digitálny omnibus z roku 2026 (nariadenie (EÚ) 2026/1744, účinné od 27. júla 2026) pridáva do článku 5 nový zákaz, ako písmená ba) a bb) článku 5 ods. 1: systémy AI na generovanie intímnych zobrazení bez súhlasu (takzvané „vyliekacie“ aplikácie) a materiálu zobrazujúceho sexuálne zneužívanie detí. Uplatňovať sa začne po prechodnom období do 2. decembra 2026, čo z neho robí jediný skutočne krátkodobý termín medzi zákazmi.

Rozdelenie zodpovednosti stojí za zapamätanie: poskytovatelia sú v ohrození, ak je takýto výstup zamýšľaným účelom *alebo* odôvodnene predvídateľným a reprodukovateľným výsledkom pri absencii záruk; nasadzujúce subjekty zodpovedajú len za úmyselné zneužitie.

STAV K AUGUSTU 2026

Text omnibusu je uverejnený, takže znenie nového zákazu je konečné: článok 5 ods. 1 písm. ba) pokrýva intímny alebo sexuálne explicitný materiál zobrazujúci identifikovateľnú osobu bez jej súhlasu, písm. bb) materiál zobrazujúci sexuálne zneužívanie detí a nový článok 5 ods. 1a hovorí, kedy je v ohrození poskytovateľ a kedy nasadzujúci subjekt. Oba sa uplatňujú od 2. decembra 2026.

Už je to zákon, stále žiadna pokuta

Tu je fakt, ktorý prekvapí každé publikum, pred ktoré ho postavím. Tieto zákazy sú vymáhateľné od 2. februára 2025. Režim sankcií, tá najvyššia úroveň 35 miliónov eur alebo 7 percent, je živý od 2. augusta 2025. A k augustu 2026 nikde v Únii neexistuje potvrdená pokuta, varovanie ani formálne konanie podľa AI Actu. (Ak uvidíte titulok o „prvých pokutách podľa AI Actu“ za náborové alebo skóringové systémy, skontrolujte zdroj: jedna taká správa kolujúca v auguste 2026 bola vymyslená a povinnosti pre vysokorizikové systémy, na ktoré sa odvoláva, sa uplatňujú až od decembra 2027.)

Prečo? Čiastočne preto, že národná mašinéria vymáhania sa ešte len skrutkuje dokopy; mnohé členské štáty meškali s určením orgánov dohľadu nad trhom. Čiastočne preto, že prvé prípady v akomkoľvek novom režime trvajú. Najbližšie k testovaciemu prípadu má rozšírenie sledovania rozpoznávaním tváří na verejné zhromaždenia v Maďarsku, o ktorom desiatky mimovládok tvrdia, že porušuje článok 5 ods. 1 písm. h), ale to je advokačný tlak, nie začaté konanie.

MÔJ POHĽAD

Ticho by som nečítal ako bezpečie. Pravidlá vás zaväzujú dnes; vymáhanie je pred nami, nie za nami. Radšej by som, aspoň nateraz, počítal s horšou možnosťou.

Zakázané, alebo len vysokorizikové?

Kým pôjdeme ďalej, dám vám intuíciu, ktorá oddeľuje túto kapitolu od nasledujúcej, lebo tá istá technológia sa stále objavuje na oboch stranách čiary. Položte tri otázky: **kto** ju používa, **na kom** a **na aký účel**.

Rozpoznávanie tváří v reálnom čase na verejnom námestí, prevádzkované políciou na skenovanie davu: zakázané, pokiaľ sa neuplatní jedna z troch úzkych výnimiek. Tie isté kamery prevádzkované supermarketom: mimo zákazu, ale vysokorizikové a pod palcom GDPR. Odvodzovanie emócií na vašich zamestnancoch: zakázané. Ten istý emočný model sledujúci únavu vodiča: povolený pod bezpečnostnou výnimkou. Bodovanie dôveryhodnosti občanov z ich spoločenského života: zakázané. Bodovanie rizika splácania žiadateľov o úver z finančných dát: povolené, ale vysokorizikové.

Zákazy zachytávajú použitia, kde sa samotný účel považuje za nezlučiteľný so základnými právami: žiadne záruky ich nenapravia. Všetko o stupienok nižšie pristane vo vysokorizikovom režime Aktu, kde odpoveď nezníe „nie“, ale „áno, s ťažkými povinnosťami“. Ten režim je ďalšia kapitola.

Overte si

1. AI mobilnej hry individualizuje bonusy, aby konkrétni hráči ďalej mŕňali, hoci vývojári nikdy ujmu nezamýšľali. Podľa článku 5 ods. 1 písm.

a) je to: A) Legálne, lebo vývojári nemali úmysel škodiť B) Potenciálne zakázané, lebo sa počíta účinko podstatného narušenia správania, nielen cieľ C) Legálne, lebo hry sú zábava a mimo rozsahu D) Len otázka transparentnosti podľa článku 50

2. Ktoré tvrdenie o zákaze sociálneho bodovania v článku 5 ods. 1 písm. c) je správne? A) Vzťahuje sa len na orgány verejnej moci B) Zakazuje všetko bodovanie jednotlivcov vrátane bankových modelov úverového rizika C) Vzťahuje sa na ktoréhokoľvek aktéra, verejného alebo súkromného, keď skóre spôsobuje škodlivé zaobchádzanie v nesúvisiacich kontextoch alebo neprimerané zaobchádzanie D) Omnibus ho odložil na december 2027

3. Polícia chce AI na predpovedanie, kto spácha trestný čin, čisto na základe osobnostného profilovania. Iný nástroj pomáha policajtom posudzovať podozrivých v existujúcom prípade postavenom na overených dôkazoch. Podľa článku 5 ods. 1 písm. d): A) Oba sú zakázané B) Ani jeden nie je zakázaný C) Prvý je zakázaný; druhý spadá pod výnimku pre podporu ľudského posúdenia založeného na objektívnych, overiteľných skutočnostiach D) Zakázaný je len druhý

4. Firma nainštaluje kamery, ktoré odvodzujú emócie zamestnancov na hodnotenie výkonu; flotila kamiónov používa kamery detegujúce únavu vodiča. Podľa článku 5 ods. 1 písm. f): A) Oba sú zakázané B) Použitie na hodnotenie výkonu je zakázané; systém na únavu sa môže oprieť o bezpečnostnú výnimku C) Oba sú povolené, ak zamestnanci súhlasia D) Pokryté sú len vzdelávacie inštitúcie, takže ani jeden nie je zakázaný

5. Nákupné centrum prevádzkuje rozpoznávanie tvárí v reálnom čase na identifikáciu vracajúcich sa zákazníkov. Ktoré posúdenie je najpresnejšie?

A) Zakázané podľa článku 5 ods. 1 písm. h), ktorý zakazuje všetku biometrickú identifikáciu v reálnom čase na verejnosti B) Úplne legálne, lebo regulované sú len vlády C) Mimo zákazu podľa článku 5 ods. 1 písm. h) (ktorý je vymedzený na presadzovanie práva), ale zachytené GDPR a klasifikované ako vysokoriziková AI D) Zakázané, pokiaľ centrum nezíska súdne povolenie

Odpovede. 1: B. Zákaz pokrýva techniky s cieľom *alebo účinkom* podstatne narušiť správanie a usmernenia Komisie čítajú „účelovo“ objektívne, takže úmysel škodiť nie je potrebný. 2: C. Zákaz je neutrálny voči aktérovi a zahryzne na dvoch vetvách ujmy, kým bežný kreditný skóring vo vlastnom kontexte je vysokorizikový, nie zakázaný. 3: C. Zákaz sa vzťahuje len na predpovede založené *výlučne* na profilovaní alebo osobnostných vlastnostiach, s výslovnou výnimkou pre AI podporujúcu ľudské posúdenie založené na faktoch. 4: B. Odvodzovanie emócií na pracovisku je zakázané, ale systémy určené zo zdravotných alebo bezpečnostných dôvodov, ako detekcia únavy, sú vyňaté. 5: C. Článok 5 ods. 1 písm. h) pokrýva diaľkovú biometrickú identifikáciu v reálnom čase len na účely presadzovania práva, takže súkromné komerčné použitie padá mimo zákazu, ale do teritória GDPR a vysokorizikového režimu.

Scenár s nákupným centrom už prezradil pointu: o stupienok pod „zakázané“ sedí oveľa väčšia kategória, vysokoriziková AI, a rozhodnúť, čo do nej patrí, je celá úloha ďalšej kapitoly.

Vysokoriziková AI: klasifikácia

Vráťte sa na chvíľu k nášmu bankovému úverovému modelu. Nakrmíte ho klientovým príjmom, platobnou históriou a pár desiatkami ďalších premenných a on predpovie, či ten človek úver splatí. Nikomu sa nič nestane. Nič nevybuchne. A predsa tento nenápadný model sedí priamo v najprísnejšie regulovanej kategórii AI Actu: vysokoriziková AI.

Prečo? Lebo Akt nemeria riziko podľa toho, aká pôsobivá je technológia. Meria ho podľa toho, čo systém môže urobiť so životom človeka. Zamietnutie úveru môže rozhodnúť, či si rodina kúpi domov. Automatizovaný triedič životopisov môže potichu ukončiť kariéru skôr, než sa začne. Otázka nikdy neznie „aká inteligentná je tá AI?“, ale „koľko je v stávke pre človeka na príjmom konci?“

Táto kapitola teda odpovedá na druhú veľkú otázku Aktu. Prvá: je to vôbec AI? Druhá: je to vysokoriziková AI? Ak áno, nasleduje dlhý zoznam povinností; to je ďalšia kapitola. Klasifikácia žije v článku 6 a prílohe III: najprv všeobecné pravidlá, potom únikový poklop, potom samotný zoznam, oblasť po oblasti.

Jeden dátum, o ktorý všetko ukotvíme: povinnosti pre vysokorizikové systémy pri prípadoch použitia z prílohy III nižšie sa uplatňujú od 2. decembra 2027, kam ich z augusta 2026 posunula novela omnibus z roku 2026. Máte čas, ale klasifikačná práca je presne ten pomalý, organizačný druh, s ktorým chcete začať skoro.

Dvoje dvere do vysokého rizika

Článok 6 vám dáva dvoje oddelené dvere do kategórie vysokého rizika. Je to ALEBO, nie A: prejsť ktorýmkoľvek stačí.

Dvere jeden: cesta cez regulované výrobky (príloha I). Predpokladajme, že vyrábate výťahy. Alebo zdravotnícke pomôcky, hračky, strojové zariadenia, tlakové zariadenia. Už sa topíte v legislatíve EÚ o výrobkoch a príloha I je jednoducho zoznam týchto existujúcich zákonov. Článok 6 ods. 1 hovorí: ak je váš systém AI bezpečnostným komponentom takéhoto výrobku, alebo ak systém AI *sám* je výrobkom, a ten výrobok

musí podľa vlastnej legislatívy prejsť posudzovaním zhody treťou stranou, potom je systém AI vysokorizikový.

Všimnite si obe polovice tej vety. Polovica o bezpečnostnom komponente je intuitívna: AI modul rozhodujúci, kedy zaberie núdzová brzda výťahu. Ale druhá polovica sa ľahko prehliadne: AI môže *byť* výrobkom. Predstavte si diagnostický softvér, ktorý analyzuje medicínske snímky. Nie je ničím bezpečnostným komponentom; je samotnou zdravotníckou pomôckou, regulovanou v rámci pravidiel pre zdravotnícke pomôcky, a do vysokého rizika vstupuje týmito istými dverami. Ak sa pýtate len „je moja AI bezpečnostný komponent?“, prehliadnete celú triedu samostatných AI výrobkov.

Dve aktualizácie omnibusu k tejto ceste. Po prvé, definícia „bezpečnostného komponentu“ sa zúžila: AI používaná čisto na pomoc používateľovi, optimalizáciu alebo kontrolu kvality je vonku, pokiaľ jej zlyhanie nemôže ohroziť zdravie alebo bezpečnosť. Po druhé, dátum: povinnosti pre túto cestu cez prílohu I sa uplatňujú od 2. augusta 2028. Ak pracujete v niektorom z týchto odvetví, moja rada je jednoduchá: otvorte prílohu I, nájdite nariadenie, ktoré vás už riadi, a pracujte odtiaľ. Tieto odvetvia majú desaťročia praxe so vstrebávaním pravidiel EÚ pre výrobky; toto je ďalšia vrstva, nie nový svet.

Dvere dva: zoznam (príloha III). Toto sú dvere, na ktorých záleží väčšine čitateľov tejto knihy, lebo zachytávajú odvetvia, ktoré nikdy predtým neboli regulované cez výrobky: banky, personálne oddelenia, univerzity, poisťovne. Príloha III je doslova zoznam prípadov použitia naprieč ôsmimi oblasťami. Ak je váš systém AI určený na použitie v jednom z nich, je vysokorizikový. Nevyžaduje sa bezpečnostný komponent, nevyžaduje sa legislatíva o výrobkoch. Známkovanie študentských esejí pomocou AI? Vysoké riziko. Bankový úverový model? Vysoké riziko. Len tým, že sú na zozname. Väčšina tejto kapitoly ten zoznam prechádza, lebo podľa mojej skúsenosti práve tu pristane takmer každá skutočná klasifikačná otázka.

Cesta von a cesta späť dnu

Ešte jeden kus článku 6 pred zoznamom: únikový poklop. Pristátie v oblasti prílohy III automaticky nerobí váš systém vysokorizikovým. To je jedna z rozumnejších dizajnových volieb Aktu. Článok 6 ods. 3 poskytuje výnimku: váš systém z prílohy III

nie je vysokorizikový, ak nepredstavuje významné riziko pre zdravie, bezpečnosť alebo základné práva, a to aj tým, že podstatne neovplyvňuje výsledok rozhodovania, a platí jedna zo štyroch podmienok.

Štyri podmienky, každá s príchutou: systém vykonáva len **úzku procedurálnu úlohu** (AI, ktorá stiahne študentské odovzdávky z úložiska a úhladne ich premenuje pre učiteľa: v oblasti vzdelávania, ale neznámkuje nič); **zlepšuje výsledok predtým dokončenej ľudskej činnosti** (AI navrhujúca stylistické vylepšenia esejí *potom*, ako učiteľ rozhodol o známke); **deteguje vzorce rozhodovania alebo odchýlky** bez toho, aby nahrádzala ľudské posúdenie; alebo vykonáva **prípravnú úlohu** pre posúdenie (plánovanie kontrolných návštev pacientov v nemocnici, bez dotyku s diagnózou).

Spoločná niť: ľudské rozhodnutie ostáva u človeka a AI ostáva na okrajoch.

MÔJ POHĽAD

Postaviť systém takto na okraje je legitímny dizajn, nie podvádžanie. Tuším, že „inžinierstvo výnimiek“ sa môže stať samostatnou konzultačnou špecializáciou.

Ale kým začnete byť kreatívni, prečítajte si vetu, ktorá prebíja všetko.

AKT HOVORÍ · ČLÁNOK 6 ODS. 3

Bez ohľadu na prvý pododsek sa systém AI uvedený v prílohe III vždy považuje za vysokorizikový, ak sa ním vykonáva profilovanie fyzických osôb.

Zrozumiteľne: ak systém profiluje ľudí, teda spracúva ich osobné údaje, aby hodnotil aspekty ich života ako výkon v práci, majetkové pomery, zdravie alebo správanie, výnimka jednoducho nie je dostupná. To na mieste zabíja dva lákavé únikové plány. Prebaliť bankový úverový model ako „čisto poradný“ druhý názor alebo ako widget na verejnom webe, ktorý vás „len informuje“, či by ste úver dostali? Stále hodnotí majetkové pomery človeka, čo je profilovanie, takže ostáva vysokorizikový bez ohľadu na to, ako poradne je zarámovaný. HR nástroj, ktorý manažérovi „len zvýrazňuje trendy“ vo výkonnostných dátach zamestnancov? Hodnotenie výkonu v práci je

učebnicové profilovanie; vysoké riziko, bodka. Výnimka je pre skutočne okrajové systémy (triediče súborov a plánovače), nie pre zmäkčené verzie jadrových bodovacích strojov.

A ak výnimku uplatníte, článok 6 ods. 4 pripája tri povinnosti, nie jednu. Musíte svoje posúdenie **zdokumentovať** pred uvedením systému na trh. Musíte systém **zaregistrovať** v databáze EÚ; áno, aj váš vlastným posúdením *ne*-vysokorizikový systém ide podľa článku 49 ods. 2 do registra; omnibus z roku 2026 zľahčil, čo musíte podať, ale povinnosť trvá. A musíte **predložiť dokumentáciu na požiadanie** vnútroštátnych príslušných orgánov. Výnimka nie je tichý odchod: necháte papierovú stopu, objavíte sa v databáze a regulátor môže požiadať, aby videl vašu prácu.

Biometria

Príloha III sa otvára biometriou, „pokiaľ je jej použitie povolené podľa príslušných právnych predpisov Únie alebo vnútroštátnych právnych predpisov“, a táto úvodná fráza je varovanie.

POZOR

Viacero položiek prílohy III sedí priamo vedľa zákazu z článku 5 a odôvodnenia sú výslovné v tom, že položky vysokého rizika pokrývajú len to, čo *nie je* už zakázané (odôvodnenie 54). Najprv skontrolujte zakázané praktiky; táto pasca sa len v tejto prvej oblasti objaví dvakrát.

Tri položky. Po prvé, diaľková biometrická identifikácia: kamerový systém skenujúci dav na vlakovej stanici voči databáze nezvestných osôb alebo podozrivých. Overenie je výslovne vylúčené. Letisková e-brána kontrolujúca, že ste osoba z vlastného pasu, nie je vysokoriziková, lebo jej jediným účelom je potvrdiť, že ste tým, za koho sa vydávate.

Po druhé, biometrická kategorizácia podľa citlivých alebo chránených atribútov. Tu je prvá pasca. Raz som sledoval konferenčnú prezentáciu o interaktívnych marketingových banneroch pre bankové pobočky: kamera nad obrazovkou číta tvár čakajúceho klienta a prispôsobuje reklamu. Odvodzovať týmto spôsobom pohlavie, vek alebo náladu je podľa tejto položky vysokorizikové. Ale vo chvíli, keď systém z tváre

odvodí *rasu alebo etnický pôvod*, už nie ste v prílohe III. Biometrická kategorizácia vyvodzujúca rasu je zakázaná podľa článku 5 ods. 1 písm. g), zakázaná od februára 2025. Tá istá kamera, iný atribút, úplne iný právny vesmír.

Po tretie, rozpoznávanie emócií, a tu je druhá pasca. Pamätáte na recepcnú luxusného hotela pod náladovou kamerou, s manažérom vybiehajúcim vždy, keď úsmev vybledne? Takéto systémy skutočne existovali. Ale odvodzovanie emócií zamestnancov na pracovisku je *zakázané* podľa článku 5 ods. 1 písm. f), nie vysokorizikové; to isté platí pre školy, s úzkou výnimkou zo zdravotných alebo bezpečnostných dôvodov. Položka rozpoznávania emócií v prílohe III pokrýva to, čo ostáva mimo zákazu: povedzme maloobchodný systém čítajúci nálady *zákazníkov*, aby prispôbil obsluhu. To je vysokorizikové, plus povinnosť informovať ľudí, ktorí sú mu vystavení.

Kritická infraštruktúra

Ďalej: AI používaná ako bezpečnostný komponent pri riadení a prevádzke kritickej digitálnej infraštruktúry, cestnej premávky alebo dodávok vody, plynu, tepla alebo elektriny. Čítajte to ako vyčerpávajúci zoznam, lebo ním je. Bankovníctvo v ňom nie je. Nech vaša banka počas pandémie pristála na akomkoľvek zozname „kritickej infraštruktúry“, s touto položkou to nemá nič spoločné, tá si berie definíciu zo smernice EÚ o kritických subjektoch. A „bezpečnostný komponent“ tu má konkrétny význam vypísaný v odôvodnení 55: niečo, čo priamo chráni fyzickú integritu alebo zdravie a bezpečnosť, bez toho, aby to bolo nevyhnutné na fungovanie systému.

AKT HOVORÍ · ODÔVODNENIE 55

Komponenty, ktoré sa majú používať výlučne na účely kybernetickej bezpečnosti, by sa nemali považovať za bezpečnostné komponenty.

Zrozumiteľne: váš AI nástroj kybernetickej obrany strážiaci cloudovú platformu *nie je* bezpečnostným komponentom podľa prílohy III, nech pôsobí akokoľvek kriticky. Čo áno? Vlastné príklady odôvodnenia: systémy riadenia požiarnych alarmov v cloudových výpočtových centrách, monitorovanie tlaku vody. Pridajte AI optimalizované semaforey na zložitej križovatke, vyvažovanie inteligentnej siete (priateľ pracujúci na optimalizácii

elektriny v Spojenom kráľovstve mi hovorí, že dopyt vyskočí cez polčas veľkých futbalových zápasov, národ zapínajúci rýchlovarné kanvice) alebo automatizované dávkovanie chemikálií v úpravni vody. Zlyhanie tam ohrozuje ľudí vo veľkom. To je test.

Vzdelávanie a odborná príprava

Štyri položky: AI rozhodujúca o prijatí do vzdelávacích inštitúcií; AI hodnotiaca výsledky vzdelávania (známkovanie esejí, projektov, skúšok s otvorenými otázkami; test s výberom odpovede s vopred danými správnymi odpoveďami nezahrňa žiadne odvodzovanie a pravdepodobne vôbec nie je AI); AI posudzujúca, k akej úrovni vzdelania má niekto prístup, vrátane prístupu k dotovaným rekvalifikačným programom; a AI dohľad pri skúškach, systémy sledujúce vašu webkameru a údery do klávesnice počas online certifikačných skúšok. Priznám sa, že toto je oblasť prílohy III, ktorú mám najradšej. Univerzity *budú* potrebovať AI nástroje na známkovanie, ako generatívna AI znásobuje študentské výstupy; Akt ich nezakazuje, len trvá na tom, aby boli postavené starostlivo.

Zamestnanosť a riadenie pracovníkov

AI na nábor a výber (triediče životopisov zoradujúce stovky uchádzačov, platformy cielenej pracovnej inzercie) a AI robíaca rozhodnutia o povýšení, prepustení, prideľovaní úloh alebo monitorovaní výkonu. Režim zlyhania je tu tichý: natrénujte náborový model na rozhodnutiach minulých manažérov, ktorí uprednostňovali jedno pohlavie alebo národnosť, a model zaujatosť poslušne zdedí. A každému náborovému manažérovi, ktorý hrdo postuje na LinkedIn o hodnotení kandidátov cez ChatGPT: prevádzkujete vysokorizikový prípad použitia AI na nástroji, ktorý halucinuje. Prosím, prestaňte. Jedna hranica z oblasti biometrie platí aj tu: monitorovanie *výkonu* je vysokorizikové; odvodzovanie *emócií* zamestnancov je zakázané.

Prístup k základným službám

Zmiešaný balík štyroch položiek, a úprimne nie moje obľúbené zoskupenie: nárok na verejné dávky, úverová bonita, poistenie a tiesňové dispečingy majú dosť odlišné

stávky. Orgány verejnej moci používajúce AI na priznanie alebo odobratie dávok ako podpora v nezamestnanosti: vysoké riziko. Hodnotenie úverovej bonity fyzických osôb, náš bankový úverový model, je vysokorizikové tiež, s jednou výnimkou: detekcia podvodov je vyňatá. Bývalí kolegovia z bankovníctva mi vysvetlili logiku: podvodníci menia taktiku každý týždeň a modely na podvody potrebujú voľné ruky na reakciu. Potom posudzovanie rizika a stanovovanie cien konkrétne v *životnom a zdravotnom* poistení (nie v poistení áut) a nakoniec AI, ktorá klasifikuje tiesňové volania alebo prioritizuje vysielanie polície, hasičov a záchraniek, vrátane triedenia pacientov.

Presadzovanie práva

Predpovedanie, kto sa môže stať obeťou trestného činu, AI polygrafy, hodnotenie spoľahlivosti dôkazov a predikcia recidívy. Táto posledná potrebuje opatrnosť, lebo sedí priamo oproti zákazu. Čítajte položku prílohy III pozorne: pokrýva posudzovanie rizika, že osoba spácha trestný čin alebo sa ho dopustí opakovane, *nie výlučne na základe profilovania*. Prečo to zvláštne znenie? Lebo článok 5 ods. 1 písm. d) *zakazuje* predikciu kriminálneho rizika založenú výlučne na profilovaní alebo osobnostných vlastnostiach: scenár Minority Report. Čo prežíva ako vysokorizikové, je užší prípad, nástroj podporujúci ľudské posúdenie už ukotvené v objektívnych, overiteľných skutočnostiach priamo spojených s trestnou činnosťou.

Aj povolená verzia nesie slávnú pascu: spätnú slučku prediktívnej policajnej práce. Pošlite hliadky tam, kde model predpovedá kriminalitu, a policajti nájdu viac kriminality presne tam, čo pretrénuje model, aby posielal ešte viac hliadok. Systém sa kŕmi sám a driftuje do zaujatosti. Rozsah je to, čo túto oblasť robí nebezpečnou.

Migrácia, azyl a kontrola hraníc

AI posudzujúca bezpečnostné, migračné alebo zdravotné riziká ľudí vstupujúcich do členského štátu; AI pomáhajúca so žiadosťami o azyl, víza a pobyt (povedzme konzulárny nástroj odhadujúci pravdepodobnosť, že žiadateľ prekročí povolený pobyt, z histórie cestovania a finančného stavu); AI polygrafy; a identifikačné systémy na hraniciach, s výnimkou overovania cestovných dokladov. Ľudské hraničné kontroly ostávajú nedotknuté; klasifikáciu spúšťa automatizácia týchto úsudkov.

Výkon spravodlivosti a demokratické procesy

Dve položky. Po prvé, AI pomáhajúca *justičnému orgánu* pri skúmaní a uplatňovaní práva: nástroje právneho rešeršu navrhujúce súdu precedensy, systémy odporúčajúce výšku trestu. Všimnite si limit. Toto pokrýva súdy, nie právneho chatbota za desať eur mesačne predávaného spotrebiteľom. Zvláštna medzera, podľa mňa, vzhľadom na to, ako seabedome generatívna AI halucinuje judikatúru.

Po druhé, AI určená na ovplyvňovanie výsledku volieb alebo referenda, s výslovnou výnimkou pre logistiku kampane v zázemí. Nosné slovo je *určená*. Nástroj mikrocílenia kampane postavený na usmerňovanie voličov do tejto položky patrí. Všeobecný odporúčací systém feedu sociálnej platformy nie: jeho zamýšľaný účel je zapojenie, nie voľby, a tú vrstvu upravuje akt o digitálnych službách a nariadenie EÚ o politickej reklame.

POZOR

Po Cambridge Analytica je lákavé čítať túto položku ako „AI Act reguluje feedy sociálnych sietí“. Nereguluje; klasifikácia jazdí na zamýšľanom účele.

Môže sa zoznam meniť?

Príloha III nie je vytesaná do kameňa, ale ruky Komisie sú menej voľné, než sa možno obávate. Podľa článku 7 môže Komisia delegovaným aktom pridávať alebo upravovať *prípady použitia*, ale len v rámci ôsmich existujúcich oblastí a len tam, kde nový prípad použitia predstavuje riziko rovnocenné alebo väčšie než tie už uvedené. Nemôže vyčarovať deviatu oblasť z ničoho; na to by bol potrebný plný legislatívny proces. Kritériá posudzovania v článku 7 ods. 2 sekajú aj na obe strany: popri faktoroch zameraných na ujmu sedí aj rozsah *prínosov* systému a to, či existujúce právo EÚ už poskytuje nápravu, čo oboje argumentuje proti pridávaniu položiek. Odstránenie je možné za vlastných dvoch kumulatívnych podmienok a Parlament aj Rada môžu proti akémukoľvek delegovanému aktu namietat.

Ešte jedna vec na váš radar: článok 6 ods. 5 vyžadoval, aby Komisia do februára 2026 zverejnila praktické klasifikačné usmernenia s vypracovanými príkladmi.

Usmernenia podľa článku 6 ods. 5 prišli neskoro a v návrhu: zverejnené 19. mája 2026, konzultácia otvorená do 23. júla 2026, konečná verzia sľúbená do konca roka. Nebol prijatý žiadny delegovaný akt podľa článku 7, omnibus z roku 2026 nechal zoznam v prílohe III nedotknutý a prvé preskúmanie Komisie podľa článku 112 ods. 1 (22. mája 2026) uzavrelo, že prílohu III zatiaľ netreba meniť.

Len zväzok návrhu k prílohe III má takmer 150 strán vypracovaných príkladov, najbližšie k oficiálnym odpovediam na otázky „je môj systém vysokorizikový?“. Berte ho ako silné orientačné usmernenie, nie ustálené právo; nestavil by som program súladu na znení, ktoré sa do decembra ešte môže posunúť.

Kde vás to necháva? S kontrolným zoznamom. Je to AI? Je v regulovanom výrobku (príloha I), ako bezpečnostný komponent alebo ako samotný výrobok? Je jeho zamýšľaný účel na zozname v prílohe III? Ak áno, platí naozaj niektorá podmienka výnimky a profiluje systém ľudí, čo tie dvere zabuchne? Zdokumentujte, zaregistrujte, buďte pripravení ukázať svoju prácu. Klasifikácia je zodpovedateľná; len odmeňuje pomalé čítanie zoznamu.

Overte si

1. Firma predáva samostatný AI softvér, ktorý analyzuje röntgenové snímky na detekciu zlomenín. Je regulovaný ako zdravotnícka pomôcka vyžadujúca posudzovanie zhody treťou stranou. Je podľa AI Actu vysokorizikový? A) Nie; nie je bezpečnostným komponentom žiadneho výrobku. B) Áno, cestou cez prílohu I, lebo systém AI je sám výrobkom. C) Len ak sa objaví aj v prílohe III. D) Nie; medicínsky softvér je z AI Actu vyňatý.

2. Hotel nainštaluje kameru, ktorá číta výrazy tváre recepčných, aby manažér mohol zasiahnuť, keď ich nálada klesne. Podľa AI Actu je tento systém: A) Vysokorizikový podľa prílohy III bodu 1 písm. c), rozpoznávanie emócií. B) S obmedzeným rizikom, vyžadujúci len oznámenie o transparentnosti. C) Zakázaný; odvodzovanie emócií zamestnancov na pracovisku je zakázané podľa článku 5 ods. 1 písm. f). D) Mimo Aktu, lebo kamera nie je otočená na zákazníkov.

3. Banka prepozicionuje svoj model schvaľovania úverov ako „čisto poradný“: konečné rozhodnutie vždy robí človek. Môže banka použiť výnimku podľa článku 6 ods. 3, aby unikla klasifikácii vysokého rizika? A) Áno; systém teraz len zlepšuje predtým dokončenú ľudskú činnosť. B) Áno; poradné systémy sú z prílohy III úplne vylúčené. C) Len ak systém najprv zaregistruje. D) Nie; hodnotenie úverovej bonity profiluje fyzické osoby a profilujúce systémy v prílohe III sú vždy vysokorizikové.

4. Poskytovateľ vlastným posúdením zaradí svoj systém z prílohy III ako ne-vysokorizikový podľa výnimky. Čo musí urobiť? A) Nič; výnimka sa uplatňuje automaticky. B) Zdokumentovať posúdenie, zaregistrovať systém v databáze EÚ a poskytnúť dokumentáciu orgánom na požiadanie. C) Získať pred spustením schválenie od notifikovanej osoby. D) Zverejniť posúdenie na svojom webe do 30 dní.

5. Čo z nasledujúceho môže Európska komisia urobiť s prílohou III delegovaným aktom podľa článku 7? A) Pridať úplne novú oblasť, napríklad „maloobchod a e-commerce“. B) Pridať alebo upraviť prípady použitia v rámci ôsmich existujúcich oblastí, kde je riziko aspoň rovnocenné už uvedeným prípadom použitia. C) Prepísať klasifikačné pravidlá v článku 6. D) Nič; príloha III sa dá zmeniť len novým nariadením.

Odpovede: 1: B. Článok 6 ods. 1 pokrýva AI, ktorá je bezpečnostným komponentom alebo je sama výrobkom podľa legislatívy z prílohy I vyžadujúcej posudzovanie zhody treťou stranou. 2: C. Odvodzovanie emócií na pracovisku je zakázanou praktikou od februára 2025; položka rozpoznávania emócií v prílohe III pokrýva len systémy, na ktoré zákaz nedosiahne, napríklad tie otočené na zákazníkov. 3: D. Posledný pododsek článku 6 ods. 3 robí profilujúce systémy vždy vysokorizikovými, takže žiadne „poradné“ prepozicionovanie na výnimku nedosiahne. 4: B. Článok 6 ods. 4 ukladá tri povinnosti: dokumentáciu pred uvedením na trh, registráciu podľa článku 49 ods. 2 a predloženie na požiadanie. 5: B. Článok 7 umožňuje pridávať alebo upravovať prípady použitia len vnútri existujúcich oblastí a len za kumulatívnych podmienok rizika; nové oblasti by potrebovali plný legislatívny proces.

Takže váš systém je na zozname, profiluje ľudí a niet cesty von: ďalšia kapitola prechádza tým, čo poskytovateľ vysokorizikového systému naozaj musí postaviť,

zdokumentovať a preukázať.

Vysokoriziková AI: povinnosti

Verdikt je teda vonku. Úverový model vašej banky, ten, ktorý rozhoduje o hypotékach, kým klient usrkáva kávu zadarmo, sedí na zozname v prílohe III, žiadna z únikových ciest z predchádzajúcej kapitoly sa neuplatní a je oficiálne vysokorizikový. Čo teraz?

Teraz prichádza časť Aktu, kde žije skutočná práca: články 8 až 15, kapitola III, oddiel 2. Sedem vecných požiadaviek: riadenie rizík, správa dát, technická dokumentácia, vedenie záznamov, transparentnosť voči nasadzujúcim subjektom, ľudský dohľad a trojica presnosť, spoľahlivosť a kybernetická bezpečnosť. V skutočných organizáciách je to medziodborový tím (dátoví vedci, právnici, pracovníci pre súlad, doménoví experti) pracujúci mesiace. Táto kapitola z vás ten tím neurobí; dá vám úprimný obraz toho, čo každý článok vyžaduje, aby ste vtedy, keď sa tím poskladá, vedeli, čo stavia a prečo.

Najprv orientácia: pri systémoch z prílohy III sa tieto povinnosti uplatňujú od 2. decembra 2027, kam ich posunula novela omnibus z roku 2026, a pri systémoch zabudovaných vo výrobkoch z prílohy I od 2. augusta 2028. Pevné kalendárne dátumy: vzdialené, ale práca je hlboká na mesiace, čo je presne dôvod, prečo boli posunuté.

Článok 9: systém riadenia rizík

Povedzme, že naša banka postaví AI náborový nástroj, vysokorizikový podľa kategórie zamestnanosti v prílohe III. Článok 9 hovorí, že systém riadenia rizík sa zavedie, implementuje, zdokumentuje a udržiava: nie kontrolný zoznam, ktorý raz vyplníte, ale nepretržitý, iteratívny proces počas celého životného cyklu systému. Plánovaná prehliadka, nie jednorazové očkovanie.

Proces má štyri opakujúce sa kroky: identifikovať a analyzovať známe a odôvodnene predvídateľné riziká pre zdravie, bezpečnosť alebo základné práva; odhadnúť a vyhodnotiť ich pri zamýšľanom použití a odôvodnene predvídateľnom nesprávnom použití; ďalej hodnotiť nové riziká z dát monitorovania po uvedení na trh, keď je systém živý; a prijať cielené opatrenia na to, čo ste našli. Článok 9 ods. 3 užitočne obmedzuje cvičenie na riziká, ktoré naozaj viete riešiť dizajnom, vývojom alebo primeranými

informáciami pre nasadzujúci subjekt. Po zmiernení musí byť zvyškové riziko posúdené ako prijateľné; čo sa nedá odstrániť inžinierstvom, zmiernite a zverejníte.

V praxi banka vymenuje medziodborový tím (zvyčajne ľudia dávajúci 10 až 20 percent svojho týždňa), ktorý vedie workshopy o tom, čo by sa mohlo pokaziť: zaujatosť pri výbere kandidátov, nesprávne prečítané životopisy, úniky súkromia. Každé riziko dostane pravdepodobnosť, závažnosť, zdokumentované zmiernenie; zraniteľné skupiny vrátane maloletých dostanú podľa článku 9 ods. 9 osobitnú starostlivosť.

Potom príde testovanie. Článok 9 ods. 8 vyžaduje testovanie v ktoromkoľvek vhodnom bode počas vývoja a v každom prípade pred uvedením na trh, voči metrikám a pravdepodobnostným prahom, ktoré ste vopred definovali. Toľko je povinné.

Testovanie v reálnych podmienkach, teda skúšanie nástroja na živých uchádzačoch o prácu, je voliteľné a regulované: článok 9 ods. 7 ho povoľuje len „v súlade s článkom 60“, čo prináša plán testovania v reálnych podmienkach, registráciu, dohľad nad trhom a podľa článku 61 informovaný súhlas ľudí, na ktorých testujete. Potichu spúšťať váš nedokončený náborový model na uchádzačoch, ktorí nikdy nesúhlasili s tým, že budú testovacími subjektmi, nie je dôslednosť; je to samo osebe porušenie.

Článok 10: údaje a správa údajov

Spomeňte pred klientmi článok 10 a niekto povie: „Správa dát? To robíme roky. GDPR, hotovo.“ Ten reflex je presne nesprávny. GDPR upravuje osobné údaje, kamkoľvek tečú. Článok 10 upravuje niečo užšie a hlbšie: tréningové, validačné a testovacie súbory dát, z ktorých sa váš model učí, lebo, ako sme si ukázali v kapitole 1, kto formuje dáta, formuje správanie.

Vezmite bankový model úverovej bonity. Článok 10 ods. 2 vyžaduje zdokumentované postupy pokrývajúce dizajnové voľby, pôvod dát a ich pôvodný účel, ich prípravu (čistenie, označovanie, obohacovanie), predpoklady, ktoré v sebe nesú, či ich je dosť a, čo je kľúčové, preskúmanie zaujatostí, ktoré by pravdepodobne poškodili zdravie, bezpečnosť alebo základné práva, plus opatrenia na ich detekciu, prevenciu a zmiernenie. Článok 10 ods. 3 dodáva, že súbory dát musia byť relevantné, dostatočne reprezentatívne a v čo najväčšej miere bez chýb a úplné vzhľadom na zamýšľaný účel. Ak bude model bodovať žiadateľov naprieč viacerými regiónmi, dáta by mali radšej

odrážať tie regióny a ich socioekonomický mix. Medzera v súlade môže existovať skôr, než sa napíše riadok kódu.

Jedno skutočne prekvapivé ustanovenie. Aby ste skontrolovali, či model diskriminuje podľa etnicity alebo zdravotného stavu, potrebujete citlivý atribút, presne tie dáta, ktorých sa bežne nesmiete dotknúť. Článok 10 ods. 5 poskytovateľom výnimočne umožňuje spracúvať osobitné kategórie osobných údajov prísne na účely detekcie a nápravy zaujatosti, pod navrstvenými zárukami: úloha musí byť nemožná s inými dátami (vrátane syntetických alebo anonymizovaných), prístup prísne kontrolovaný a zdokumentovaný, žiadny ďalší prenos, vymazanie po náprave zaujatosti. Akt radšej chce, aby ste s citlivými dátami zaobchádzali opatrne, než aby ste ostali nevedomí o diskriminácii svojho modelu.

Článok 11: technická dokumentácia

Článok 11 znie ako ten nudný a je čokoľvek, len nie voliteľný. Technická dokumentácia sa musí vypracovať *pred* uvedením systému na trh, nie rekonštruovať dodatočne, a musí sa udržiavať aktuálna. Jej úloha: preukázať vnútroštátnemu orgánu alebo notifikovanej osobe, ktorá ju číta bez kontextu, že systém spĺňa všetko v tomto oddiele.

Osnovu si nevymýšľate: článok 11 ods. 1 hovorí, že dokumentácia obsahuje aspoň prvky z prílohy IV, deväť z nich:

1. Všeobecný opis: zamýšľaný účel, poskytovateľ, verzovanie, interakcie s iným hardvérom a softvérom, formy, v ktorých sa dodáva (API, zabudovaný, na stiahnutie), hardvér, na ktorom beží, rozhranie nasadzujúceho subjektu a návod.
2. Podrobný opis systému a jeho vývoja: metódy, predtrénované komponenty tretích strán, dizajnové voľby a kompromisy, architektúra, pôvod dát, posúdenie ľudského dohľadu, validácia a testovanie s podpísanými testovacími správami, opatrenia kybernetickej bezpečnosti.
3. Podrobné informácie o monitorovaní, fungovaní a kontrole: schopnosti a limity, presnosť pre konkrétne skupiny, predvídateľné nezamýšľané výsledky, špecifikácie vstupov.
4. Prečo sú vami zvolené výkonnostné metriky pre tento systém vhodné.

5. Podrobný opis systému riadenia rizík podľa článku 9.
6. Relevantné zmeny vykonané počas životného cyklu.
7. Zoznam uplatnených harmonizovaných noriem alebo, kde žiadne nie sú, opis riešení, ktoré ste namiesto nich prijali (k polovici roka 2026 sú to všetci; viac o tom na konci).
8. Kópia EÚ vyhlásenia o zhode (článok 47).
9. Opis systému a plánu monitorovania po uvedení na trh (článok 72).

Dve úľavy: obsah sa uplatňuje „podľa potreby“, primerane systému, a MSP vrátane startupov môžu použiť zjednodušený formulár, ktorý má poskytnúť Komisia a ktorý notifikované osoby musia akceptovať.

Článok 12: vedenie záznamov

Článok 12 žiada, aby vysokorizikové systémy technicky umožňovali automatické zaznamenávanie udalostí (logov) počas svojej životnosti. Predstavte si mestský AI regulátor semaforov: každá zmena signálu, každá anomália, každý alert na ľudský zásah, s časovou pečiatkou, aby sa porucha o tretej ráno dala zrekonštruovať: čo systém videl, čo rozhodol?

Ale článok je odstupňovaný. Pri väčšine vysokorizikových systémov článok 12 ods. 2 vyžaduje logovanie *primerané zamýšľanému účelu*: dosť na to, aby sa dali identifikovať situácie, keď systém môže predstavovať riziko alebo bol podstatne zmenený, aby sa nakrmilo monitorovanie po uvedení na trh a aby nasadzujúce subjekty mohli monitorovať prevádzku. Čo „dosť“ znamená pre regulátor semaforov, sa líši od úverového modelu. Vy to navrhnete, vy to zdôvodníte.

Slávny podrobný zoznam (začiatok a koniec každého použitia, referenčná databáza, voči ktorej sa kontrolovalo, vstupné dáta, ktoré viedli k zhode, identita ľudí, ktorí výsledok overili) je článok 12 ods. 3 a vzťahuje sa len na systémy diaľkovej biometrickej identifikácie podľa prílohy III bodu 1 písm. a).

POZOR

„Referenčná databáza“ tu znamená databázu tváří, voči ktorej systém páruje, nie váš dátový sklad. A ak brožúra dodávateľa tvrdí, že každý vysokorizikový systém musí logovať tie štyri položky, dodávateľ sa mýli.

Článok 13: transparentnosť voči nasadzujúcim subjektom

Situácia, ktorú Akt úhladne predvída: ste poskytovateľ predávajúci svoj vysokorizikový systém úverového skóringu iným bankám, v slovníku Aktu nasadzujúcim subjektom. Vaša dokumentácia podľa článku 11 obsahuje všetko: architektúru, tréningové dáta, obchodné tajomstvá, know-how. Odovzdajte ju zákazníkovi a odovzdali ste ju konkurentovi.

Článok 13 je odpoveď: druhý, navonok smerovaný dokument, návod na použitie, ktorý musí byť stručný, úplný, správny a jasný. Minimálne: kto ste; zamýšľaný účel systému, jeho schopnosti a obmedzenia; metriky presnosti, voči ktorým bol testovaný, jeho spoľahlivosť a kybernetická bezpečnosť a čo by ich mohlo zhoršiť; predvídateľné riziká a nesprávne použitie; kde je to relevantné, výkon na konkrétnych skupinách; vstupné dáta, ktoré potrebuje; ako interpretovať výstup („skóre nad 70 typicky znamená schválenie“); opatrenia ľudského dohľadu; výpočtové zdroje a očakávaná životnosť; a ako môže nasadzujúci subjekt zbierať a čítať logy. Poskytovateľ teda udržiava dva artefakty: hrubý interný spis pre regulátora a poctivý používateľský manuál pre zákazníka, od ktorého článok 26 očakáva, že ho nasadzujúci subjekt naozaj použije.

Článok 14: ľudský dohľad

Článok 14 vyžaduje, aby vysokorizikové systémy boli navrhnuté tak, aby ich fyzické osoby mohli počas používania účinne dozerať, s cieľom predchádzať rizikám pre zdravie, bezpečnosť a základné práva alebo ich minimalizovať. Nie dohľad ako vyhlásenie politiky: dohľad navrhnutý do produktu, cez samotné rozhranie človek – stroj.

Článok 14 ods. 4 vypisuje, čo dozerajúci človek musí byť schopný robiť: rozumieť schopnostiam a limitom systému; ostať si vedomý **automatizačného skreslenia**

(automation bias), našej dobre zdokumentovanej tendencie prehnané dôverovať sebavedomému stroju; správne interpretovať výstup; rozhodnúť sa systém nepoužiť alebo prepísať jeho výsledok; a zasiahnuť alebo ho zastaviť tlačidlom stop či podobným postupom. Niektoré opatrenia zabuduje poskytovateľ, iné identifikuje pre nasadzujúci subjekt na implementáciu (článok 14 ods. 3): zdieľaný dizajnový problém.

Dva náčrty. MediScan AI navrhuje diagnózy z röntgenov so skóre istoty; každý návrh posudzuje rádiológ, prípady s nízkou istotou idú k druhému rádiológovi a funkcia stop vráti manuálnu analýzu, ak systém začne produkovať nezmysly. LoanDecider AI boduje žiadosti od 0 do 100; stredné skóre idú k úverovému pracovníkovi, závažné zamietnutia upozornia seniorného a každé prepísanie sa loguje. Rôzne domény, rovnaká gramatika: prahy, eskalácia, školenie proti automatizačnému skresleniu, prepísanie, audítorská stopa.

Potom je tu biometrický osobitný prípad, kde Akt prestane navrhovať a začne nariaďovať.

AKT HOVORÍ · ČLÁNOK 14 ODS. 5

V prípade vysokorizikových systémov AI uvedených v prílohe III bode 1 písm. a) musia byť opatrenia podľa odseku 3 tohto článku také, aby sa navyše zabezpečilo, že nasadzujúci subjekt na základe identifikácie vyplývajúcej zo systému nekoná ani neprijme žiadne rozhodnutie, pokiaľ takáto identifikácia nebola zvlášť overená a potvrdená aspoň dvoma fyzickými osobami, ktoré majú potrebnú spôsobilosť, odbornú prípravu a právomoc.

Dvaja ľudia, zo zákona, s výnimkou tam, kde právo Únie alebo vnútroštátne právo považuje túto požiadavku za neprimeranú v kontextoch presadzovania práva, migrácie, kontroly hraníc alebo azylu. Zákonné pravidlo, nie dizajnový návrh.

A buďte opatrní so scenárom, ktorý si tu ľudia predstavia najochotnejšie: polícia skenujúca živé kamerové záznamy vo verejných priestoroch a hľadajúca hľadané osoby. To väčšinou nie je problém dizajnu dohľadu: diaľková biometrická identifikácia v reálnom čase vo verejne prístupných priestoroch na účely presadzovania práva je *zakázaná* praktika podľa článku 5 ods. 1 písm. h), zákonná len pre tri úzke ciele, len

tam, kde to vnútroštátne právo umožňuje, a len s predchádzajúcim súdnym alebo nezávislým povolením. Dobrý dizajn dohľadu nevie vyprať zakázanú prax na povolenú. Pravidlo dvoch osôb žije v zákonnom biometrickom teritóriu, napríklad pri následnej identifikácii voči záznamom závažného trestného činu.

Článok 15: presnosť, spoľahlivosť, kybernetická bezpečnosť

Článok 15 je krátky a čitateľný. Je aj, tvrdil by som, najdrahším článkom oddielu: každé z jeho troch titulných slov skrýva inžiniersku disciplínu.

Presnosť. Systém musí dosahovať „primeranú úroveň presnosti“ a fungovať konzistentne počas celého životného cyklu, s úrovňami a metrikami deklarovanými v návode na použitie (článok 15 ods. 3). Kto rozhodne, čo znamená „primeraná“? Väčšinou vy. To nie je diera; to je dizajn. Neexistujú žiadne povinné referenčné hodnoty presnosti na úrovni EÚ: článok 15 ods. 2 len poveruje Komisii, aby *podporovala* referenčné hodnoty a metodiky merania, a k augustu 2026 žiadne neexistujú. Pri väčšine systémov z prílohy III vás článok 43 ods. 2 pošle cestou vnútornej kontroly podľa prílohy VI: posudzujete vlastnú zhodu a musíte ju vedieť obhájiť. Deklarujte metriku, zdôvodnite ju, splňte ju, plňte ju ďalej.

Výber metriky je tam, kde zahryzne doménový úsudok. Vezmite model, ktorý z videa chôdze človeka skrútuje Parkinsonovu chorobu. Jeho titulná „presnosť“, podiel všetkých správnych predpovedí, môže byť takmer bezcenné číslo. Katastrofická chyba je falošne negatívny výsledok: povedať niekomu, kto chorobu má, že je v poriadku, takže nikdy nepôjde k lekárovi. Falošne pozitívny výsledok len pošle zdravého človeka na vyšetrenie, ktoré ho upokojí. Pri tom systéme vám záleží na citlivosti (recall): z ľudí, ktorí chorobu naozaj majú, aký podiel zachytíte? Iný systém vyžaduje inú voľbu, a preto vás bod 4 prílohy IV núti zdôvodniť metriky písomne.

Potom je tu „konzistentne počas celého životného cyklu“. Modely sú pri nasadení zmrazené; svet nie. Vzorce driftujú, ako inflácia pomaly zneplatňujúca pravidlo naučené na minuloročných príjmoch, alebo sa lámu, ako COVID-19 cez noc zlomil spotrebiteľské správanie. Súlad tu znamená monitorovať drift a pretrénovať na čerstvých dátach skôr, než výkon klesne pod to, čo ste deklarovali.

Spôľahlivosť. Systém musí byť čo najodolnejší voči chybám, poruchám a nezrovnalostiam, cez technické a organizačné opatrenia: redundanciu, zálohy, plány odolné voči zlyhaniu.

MÔJ POHĽAD

Toto je požiadavka, ktorá bude bolieť kultúrne. Dátoví vedci boli desaťročie školení stavať systémy, ktoré sa dobre *učia*; robiť ich tak, aby dobre *zlyhávali*, bola práca niekoho iného. Očakávajte, že účet za spoľahlivosť sa zaplatí v inžinierskych praktikách (starostlivé pipeline príznakov, spracovanie poškodených vstupov, záťažové testy), ktoré mnohé AI tímy nikdy nemuseli ovládať.

Kybernetická bezpečnosť. Článok 15 ods. 5 je najkonkrétnejší: systém musí odolávať pokusom neoprávnených tretích strán zmeniť jeho používanie, výstupy alebo výkon, s opatreniami, kde je to vhodné, proti štyrom pomenovaným útokom špecifickým pre AI. Otrávenie dát je poškodenie tréningových dát: trik z kapitoly 1, kde falošní „klienti“ krmili bankový úverový model vymyslenými záznamami, kým sa jeho rozhodovacia hranica neposunula. Otrávenie modelu mieri tou istou myšlienkou na predtrénované komponenty: open-source model, ktorý ste priskrutkovali do svojej pipeline, dorazí s cudzou manipuláciou zabezčenou vnútri. Adverzariálne príklady sú vstupy skonštruované tak, aby model urobil chybu. Útoky na dôvernosť model nahovoria, aby vyzradil svoje tréningové dáta alebo vnútornosti. Klasická IT bezpečnosť stále platí; tento zoznam je to, čo je nové teraz, keď sa softvér učí.

Praktické dôsledky

Odstúpte a tvar je jasný: články 8 až 15 opisujú zrelú inžiniersku organizáciu, ktorá pozná svoje riziká, svoje dáta, svoj výkon a vie to všetko preukázať na papieri, pred nasadením aj po ňom. Článok 8 to celé kalibruje: súlad sa posudzuje voči zamýšľanému účelu systému a aktuálnemu stavu vývoja.

Obava, ktorú od tímov počúvam neustále: „Ak sa náš prototyp počíta ako vysokoriziková AI, nepotrebuje teraz každý experiment dokumentáciu, rizikové spisy,

schválenie nejakého regulátora? To by strojnásobilo cenu skúšania nápadov.“ Akt na to odpovedá priamo.

AKT HOVORÍ · ČLÁNOK 2 ODS. 8

Toto nariadenie sa neuplatňuje na žiadnu výskumnú, testovaciu ani vývojovú činnosť týkajúcu sa systémov AI alebo modelov AI pred ich uvedením na trh alebo do prevádzky. Takéto činnosti sa vykonávajú v súlade s uplatniteľným právom Únie. Toto vylúčenie sa nevzťahuje na testovanie v reálnych podmienkach.

Zrozumiteľne: povinnosti sa viažu na uvedenie systému na trh alebo do prevádzky, nie na kútenie v laboratóriu. Váš backlog surových nápadov, vaše trojmesačné prototypy, deväť z desiatich experimentov, ktoré potichu zomrú: nič z toho AI Act nereguluje. Žiadna povinnosť nikam predkladať prototypy; regulačné experimentálne prostredia (články 57 až 59) sú dobrovoľná podpora, nie brána.

Všimnite si však poslednú vetu: vo chvíli, keď vaše „testovanie“ zahŕňa skutočných ľudí, skúšanie náborového modelu na skutočných uchádzačoch, ste späť v režime článku 60 z článku 9.

A podrobné „ako“ sa stále formuje.

STAV K AUGUSTU 2026

V Úradnom vestníku bolo citovaných nula harmonizovaných noriem podľa AI Actu, takže nikto si zatiaľ nemôže nárokovať predpoklad zhody, ktorý dodržiavanie normy raz prinesie; všetci dokumentujú vlastné riešenia podľa bodu 7 prílohy IV. Certifikát ISO/IEC 42001, nech konzultant, ktorý ho predáva, hovorí čokoľvek, je užitočné lešenie, nie súlad s AI Actom, a nedáva žiadny právny predpoklad.

Každý odznak „certifikované podľa AI Actu“ berte s veľkou rezervou a svoj konkrétny prípad prekonzultujte s právnym poradcom. Útechou je kalendár: december 2027 je dosť ďaleko na to, aby ste to urobili poriadne.

Overte si

1. Váš tím chce postaviť rýchly interný prototyp AI nápadu, ktorý je kandidátom na vysoké riziko. Povinnosti podľa článkov 8 až 15: A) Platia od prvého riadku kódu, keďže prípad použitia je vysokorizikový B) Platia až pri uvedení na trh alebo do prevádzky; výskum a vývoj pred uvedením na trh je vylúčený (článok 2 ods. 8) C) Platia až po tom, ako národné experimentálne prostredie prototyp schváli D) Nikdy neplatia pre systémy vyvinuté interne

2. Banka chce testovať svoj nevydaný AI náborový nástroj na skutočných uchádzačoch o prácu. Podľa Aktu je to: A) Povinné podľa článku 9 pred uvedením na trh B) Voľne povolené, keďže systém ešte nie je na trhu C) Voliteľné a zákonné len v režime testovania v reálnych podmienkach podľa článkov 60/61, vrátane informovaného súhlasu účastníkov D) Zakázané za každých okolností

3. Podrobný zoznam logovania v článku 12 ods. 3 (obdobie každého použitia, referenčná databáza, vstupné dáta zhody, identita overujúcich osôb) sa vzťahuje na: A) Všetky vysokorizikové systémy AI bez výnimky B) Len systémy diaľkovej biometrickej identifikácie podľa prílohy III bodu 1 písm. a); ostatné potrebujú logovanie primerané účelu C) Len systémy používané orgánmi verejnej moci D) Len systémy, ktoré sa po nasadení ďalej učia

4. Ktoré tvrdenie o presnosti podľa článku 15 je k roku 2026 správne? A) EÚ zverejnila povinné minimálne referenčné hodnoty presnosti pre každý prípad použitia z prílohy III B) Požadovanú úroveň presnosti pre každý vysokorizikový systém stanovuje notifikovaná osoba C) Na presnosti záleží len v okamihu uvedenia na trh, nie potom D) Poskytovateľ deklaruje primeranú úroveň presnosti a metriky v návode na použitie; Komisia iba podporuje vývoj referenčných hodnôt a väčšina systémov z prílohy III sa posudzuje sama

5. Kým nasadzujúci subjekt koná na základe zhody zo systému biometrickej identifikácie podľa prílohy III bodu 1 písm. a), identifikácia musí byť spravidla: A) Potvrdená zákazníckou podporou poskytovateľa B) Zvlášť overená a potvrdená aspoň dvoma kompetentnými fyzickými osobami C) Zalogovaná, bez potreby ľudského potvrdenia D) Trikrát znovu spustená samotným systémom

Odpovede: 1: B. Článok 2 ods. 8 vylučuje z rozsahu výskum, testovanie a vývoj pred uvedením na trh; experimentálne prostredia sú dobrovoľné. 2: C. Článok 9 ods. 7 hovorí, že testovanie *môže* zahŕňať testovanie v reálnych podmienkach „v súlade s článkom 60“, čo prináša plán, registráciu, dohľad a informovaný súhlas podľa článku 61. 3: B. Článok 12 ods. 3 je výslovne obmedzený na biometrickú identifikáciu podľa prílohy III bodu 1 písm. a); všetci ostatní dostávajú logovanie primerané účelu podľa článku 12 ods. 2. 4: D. Článok 15 ods. 2 iba žiada Komisiu, aby podporovala vývoj referenčných hodnôt, článok 15 ods. 3 necháva poskytovateľa deklarovať metriky a článok 43 ods. 2 smeruje väčšinu systémov z prílohy III na vlastné posúdenie. 5: B. Článok 14 ods. 5 vyžaduje samostatné overenie aspoň dvoma kompetentnými fyzickými osobami, s úzkou výnimkou tam, kde to právo považuje za neprimerané.

To boli ťažké povinnosti pre vysokorizikovú menšinu; ďalej prichádzajú povinnosti v oblasti transparentnosti, ktoré dosahujú na oveľa viac systémov, vrátane chatbota, s ktorým ste sa dnes ráno rozprávali.

Povinnosti v oblasti transparentnosti

Sú dve v noci a píšete do okna zákazníckej podpory: „Dobrý deň, som Anna, ako vám dnes môžem pomôcť?“ Anna odpovedá okamžite, v celých vetách, s veselosťou, akú o druhej ráno nemá žiadny človek. Je Anna človek? Tušíte, že nie. Ale tušiť nie je vedieť a firma má každý komerčný podnet nechať vás hádať, lebo výskum stále naznačuje, že ľudia si prácu vyprodukovanú AI cenia menej než tú ľudskú. Keby som sa v tomto kurze nahradil digitálnym avatarom a priznal to, pravdepodobne by ste kurz už len z toho dôvodu hodnotili horšie.

Kapitola IV AI Actu existuje, aby tú medzeru zatvorila: medzeru medzi tým, čím AI je, a tým, čo vám o nej povedia. Je to najkratšia kapitola celého nariadenia: jeden článok, článok 50. Ale nenechajte sa veľkosťou oklamať. Kým mašineria vysokého rizika z kapitoly III sa vzťahuje na relatívne úzky zoznam systémov, článok 50 sa dotýka takmer každej firmy, ktorá prevádzkuje chatbota, generuje obrázky alebo publikuje AI napísaný text. Pre mnohých čitateľov tejto knihy je táto kapitola časť AI Actu, ktorá vám naozaj pristane na stole ako prvá.

A už pristála. Novela digitálny omnibus z roku 2026 (nariadenie (EÚ) 2026/1744, účinné od 27. júla 2026) posunula veľké termíny pre vysokorizikové systémy na december 2027 a august 2028. Článok 50 **neposunula**. Povinnosti v oblasti transparentnosti sa uplatňujú od **2. augusta 2026**, presne ako bolo pôvodne naplánované.

POZOR

Ak ste celý AI Act založili pod „to je teraz problém roku 2027“, táto kapitola je korekcia. Omnibus nechal článok 50 na pokoji: pri transparentnosti je súlad problémom dneška, nie budúceho roka.

Článok 50 obsahuje štyri odlišné povinnosti. Najľahší spôsob, ako ich udržať oddelene, je položiť pri každej dve otázky: *kto* povinnosť nesie (poskytovateľ, ktorý systém

postavil, alebo nasadzujúci subjekt, ktorý ho používa) a čo sa musí zverejniť. Najprv vezmeme povinnosti smerom k ľuďom, potom povinnosti týkajúce sa obsahu.

Povedzte ľuďom, že hovoria so strojom

Článok 50 ods. 1 je pravidlo Anny. Poskytovatelia musia zabezpečiť, aby systémy AI určené na priamu interakciu s fyzickými osobami boli navrhnuté tak, aby tieto osoby boli informované, že majú do činenia so systémom AI.

Všimnite si zúženie. Nepokrýva to každý systém AI, ktorý prevádzkujete, len tie určené na *priamu* interakciu s ľuďmi. Váš chatbot pre zákazníkov: áno. Model kreditného skóringu z nášho bankového príkladu, potichu rozhodujúci o žiadostiach o úver v zázemí: nie. Zákazníci sa s ním nikdy nerozprávajú, takže niet čo oznamovať. (Bankový model má vlastné, oveľa ťažšie povinnosti; to boli posledné dve kapitoly.)

Toto pravidlo ste v divočine videli roky: malá ikonka robota, úvod „Dobrý deň, som virtuálny asistent“. Od augusta 2026 sa táto zdvorilosť stáva zákonnou požiadavkou. Musím povedať, že toto pravidlo sa mi páči. Vzhľadom na podnet vydávať AI za človeka je priame „musíte to povedať“ presne ten druh pravidla, ktoré poctivé firmy nič nestojí a tie ostatné chytí.

Existujú dva únikové poklapy a oba sú rozumné. Po prvé, žiadne zverejnenie netreba tam, kde je to zrejmé „z hľadiska fyzickej osoby, ktorá je primerane informovaná, pozorná a obozretná“. Ak zákazník vojde do vašej pobočky a rozpráva sa s viditeľne fyzickým robotom, nepotrebuje tabuľu „tento robot je robot“. Po druhé, povinnosť sa nevzťahuje na systémy AI povolené zákonom na odhaľovanie, vyšetrovanie alebo stíhanie trestných činov, tá istá výnimka pre presadzovanie práva, akú sme videli v kapitole o zakázaných praktikách. Predstavte si políciu používajúcu syntetický hlas pri utajenej operácii; nezastaví sa, aby to zverejnila.

Keď vás systém číta: rozpoznávanie emócií a biometrická kategorizácia

Druhá povinnosť smerom k ľuďom je tá, o ktorej takmer nikto nehovorí, čo je presne dôvod, prečo chceme, aby ste si ju zapamätali. Článok 50 ods. 3: nasadzujúce subjekty

systemu na rozpoznávanie emócií alebo systému biometrickej kategorizácie musia informovať fyzické osoby, ktoré sú mu vystavené, že systém je v prevádzke, a musia príslušné osobné údaje spracúvať v súlade s GDPR a jeho súrodencami pre presadzovanie práva.

Všimnite si, kto túto nesie: *nasadzujúci subjekt*, organizácia používajúca systém, nie firma, ktorá ho postavila. A všimnite si spúšťač: samotné *vystavenie*. Osoba nemusí s ničím interagovať; ak ju kamera číta, musí jej to byť povedané.

Možno si pamätáte našu recepcnú pod náladovou kamerou z kapitoly o zakázaných praktikách, zamestnávateľa odvodzujúceho emócie osoby v práci. Presne ten scenár nekončí tu; končí priamo zakázaný podľa článku 5, lebo rozpoznávanie emócií na pracovisku a vo vzdelávaní je zakázané. Článok 50 ods. 3 upravuje to, čo ostáva: zákonné použitia. Maloobchodník prevádzkujúci emočnú analýzu na zákazníkoch prezerajúcich si obchod, call centrum, ktorého softvér odhaduje úroveň frustrácie volajúceho, systém triediaci ľudí do biometrických kategórií na podujatí: systém môže smieť bežať, ale ľudia, ktorí sú mu vystavení, majú právo vedieť, že beží. Potichu čítať tváre zákazníkov bez slova prestáva byť v auguste 2026 rastovým trikom a stáva sa porušením.

Platí obvyklá výnimka pre presadzovanie práva, kde to zákon povoľuje a so zárukami. Pre všetkých ostatných: ak váš systém číta ľudí, povedzte to ľuďom.

Syntetický obsah verus deep fake

Zvyšné dve povinnosti sa týkajú obsahu, ktorý AI produkuje, a kým bude dávať zmysel ktorákolvek z nich, potrebujeme si ujasniť dva pojmy: *syntetický obsah* a *deep fake*.

Syntetický obsah je široká kategória: čokoľvek vygenerované AI, či už zvuk, obraz, video alebo text. Je lákavé povedať „čokoľvek, čo prejde cez AI, je syntetický obsah“, a ak čítate len odôvodnenie 133, ktoré široko hovorí o systémoch AI generujúcich „veľké množstvá syntetického obsahu“, mohli by ste tam pristáť. Ale samotný Akt kreslí hranicu, priamo v článku 50 ods. 2, ustanovení o označovaní, ktoré o chvíľu stretneme: povinnosť sa neuplatňuje tam, kde AI vykonáva „pomocnú funkciu pri štandardných redakčných úpravách“ alebo „podstatne nemení vstupné údaje poskytnuté nasadzujúcim subjektom alebo ich sémantiku“.

Takže: napíšete si vlastnú správu a požiadate ChatGPT, aby ju skorigoval. Gramatika, preklepy, tu a tam hladšia veta. Význam je váš; AI asistovala. Ten výstup označovanie nepotrebuje. Napíšete „napíš mi esej o nebezpečenstvách AI“ a výsledok skopírujete: to je syntetický obsah a povinnosť označovania sa uplatní. Čiara je v tom, či AI podstatne zmenila váš vstup alebo jeho význam, nie v tom, či bola nejaká AI niekde v miestnosti. Návrh usmernení Komisie posúva hranicu ďalej von aj na druhej strane: zoznamy odporúčaní, vygenerovaný zdrojový kód, krátke sekvencie čísel, výstupy stroj – stroj. Nič z toho sa na účely povinnosti označovania nepočíta ako syntetický obsah.

To odpovedá aj na dlhotrvajúcu debatu o Photoshope. AI odstraňovanie škvŕn a korekcia farieb na vašej vlastnej fotke sú pomocné štandardné úpravy. Vypromptovať celú scénu z ničoho je generovanie. Oba konce sú jasné; stred udrží právnikov zamestnaných.

Druhý pojem Akt definuje priamo a definícia má zvrät.

AKT HOVORÍ · ČLÁNOK 3 BOD 60

„deep fake“ je obrazový obsah, audioobsah alebo videoobsah vytvorený alebo zmanipulovaný pomocou AI, ktorý sa podobá na existujúce osoby, objekty, miesta, subjekty alebo udalosti a osobe by sa falošne javil ako autentický alebo pravdivý;

Zrozumiteľne: deep fake je podmnožina syntetického obsahu, ktorá napodobňuje niečo *skutočné* (existujúcu osobu, miesto alebo udalosť) dosť presvedčivo na to, aby to niekto mohol považovať za pravé. Musia platiť obe polovice.

MÔJ POHĽAD

Tá druhá polovica ma skutočne znepokojuje: dáva obhajobu presne nesprávnym ľuďom. „Áno, vygenerovali sme to, ale pozrite, aké je to nepresvedčivé. Väčšina divákov to preukla, takže právne to ani nie je deep fake.“ Prečo dávať zlým aktérom tento argument? Nemám dobrú odpoveď; definícia je taká, aká je.

Držte si intuitívny obraz: syntetický obsah je veľký kruh, deep fake menší kruh vnútri neho. Veľký kruh sa musí *označiť*, aby ho stroje vedeli rozpoznať; ten menší sa musí

navyše *zverejniť* ľuďom, ktorí sa naň pozerajú. V tomto poradí.

Označovanie syntetického obsahu

Teraz povinnosť, kde sa pôda naposledy pohla najviac. Článok 50 ods. 2 mieri na poskytovateľov systémov AI, výslovne vrátane AI na všeobecné účely, ktoré generujú syntetický zvukový, obrazový, video alebo textový obsah. Myslite OpenAI, myslite Midjourney. Tu je operatívna veta, lebo na presnom znení záleží.

AKT HOVORÍ · ČLÁNOK 50 ODS. 2

Poskytovatelia systémov AI vrátane systémov AI na všeobecné účely, ktoré generujú syntetický zvukový, obrazový, video alebo textový obsah, zabezpečia, aby boli výstupy systému AI označené v strojovo čitateľnom formáte a zistiteľné ako umelo vygenerované alebo zmanipulované.

Zrozumiteľne: iný stroj musí byť schopný pozrieť sa na obsah a povedať, že ho vyrobila AI. Všimnite si, že sú to *dve* kumulatívne požiadavky, nie jedna: výstup musí niesť strojovo čitateľné označenie **a** musí existovať fungujúci spôsob, ako ho zistiť. Návrh usmernení Komisie je výslovný v tom, že vodoznak, ktorý nikto nevie naozaj prečítať, sa nepočíta. Označenie bez zistiteľnosti nebude stačiť. Hranica z predchádzajúcej časti robí vrátnika: skorigovaná správa ostáva vonku, esej od nuly je vnútri.

Ako obsah vlastne označiť? Keď sa Akt písal, boli to otvorené špekulácie. Vodoznaky? Metadáta? Nové formáty súborov? Text bol zjavný bolehlav: neviditeľný vodoznak prežije vnútri obrazového súboru, ale text sa kopíruje a vkladá a viditeľná pätička „napísal ChatGPT“ sa zmaže jedným úderom klávesu. To „ako“ konečne dostalo oficiálny tvar, v dvoch dokumentoch.

Komisia prijala konečné usmernenia k povinnostiam podľa článku 50 20. júla 2026 (návrh pochádzal z mája). Kódex postupov pre transparentnosť obsahu vygenerovaného AI, zverejnený 10. júna 2026, Komisia posúdila ako primeraný 8. júla a rada pre AI 9. júla, s vyše 180 signatármi do konca júla. Spolu s ním Komisia zverejnila bezplatné ikony EÚ na označovanie AI obsahu („AI“, „AI generated“, „AI modified“), ktoré môže použiť ktokoľvek. Kým sa na ktorýkoľvek dokument spoľahnete, skontrolujte aktuálne verzie; usmernenia sa revidujú.

Je to ten istý mechanizmus, ktorý stretnete pri AI na všeobecné účely v ďalšej kapitole: podpíšte kódex, dodržiavajte ho a máte silný praktický argument, že ste v súlade. Ak poskytujete generatívny systém, tieto dva dokumenty sú váš zoznam čítania; neviditeľné vodoznaky a vložené metadáta sú presne to, o čom hovorí.

Dve spresnenia z usmernení idú proti intuíciam, ktoré možno máte.

Po prvé, formulácia „pokiaľ je to technicky možné“ v článku 50 ods. 2, časť o nákladoch na implementáciu a aktuálnom stave vývoja, podmieňuje *kvalitu* vášho technického riešenia, nie to, či musíte označovať. Nemôžete uvažovať „moje obrázky sú štylizované fantasy umenie, nie realistické, takže pravidlá označovania sa na mňa vzťahujú menej“. V článku 50 ods. 2 nie je žiadny gradient realizmu; na realizme záleží pri oddelenej otázke deep fake, ktorá príde nižšie. Usmernenia umožňujú úzke, kumulatívne výnimky z dôvodov proporcionality (uzavreté zabudované prostredia ako navigácia vo vozidle alebo prísne technické B2B výstupy ako inžinierske návrhy), ale obsah pre verejnosť sa cez tie dvere nedostane.

Po druhé, ak ste sa pýtali, či budú platformy sociálnych médií nútené skenovať a označovať každý nahraný obsah: podľa článku 50 nie. Platforma, ktorá obsah iba šíri, bez právomoci nad systémom AI, ktorý ho vyrobil, tu nie je nasadzujúcim subjektom. Usmernenia platformy iba *povzbudzujú*, aby označenie zachovali. Povinnosti označovania na strane platformy, kde existujú, žijú v inej legislatíve, nie v tomto článku.

Jeden dátumový záhyb, s láskavým dovolením omnibusu: povinnosť označovania dostáva odklad do **2. decembra 2026**, ale len pre generatívne systémy, ktoré už boli

na trhu k 2. augustu 2026. Uved'ite nový generatívny systém na trh po tomto dátume a musíte byť v súlade od prvého dňa. Je to štvormesačné vyjdenie v ústrety existujúcim systémom, nie odloženie pravidla.

Deep fake: povolené, ale označené

Nakoniec povinnosť, o ktorej počul každý. Článok 50 ods. 4 hovorí, že subjekty nasadzujúce systém AI, ktorý generuje alebo manipuluje obrazový, zvukový alebo video obsah predstavujúci deep fake, musia zverejniť, že obsah bol umelo vytvorený alebo zmanipulovaný. Všimnite si, čo tam nie je: nehovorí sa, že deep fake sú zakázané. Keď som AI Act čítal prvýkrát, bolo to moje najväčšie prekvapenie: deep fake sú v základe *povolené*; len ich musíte označiť. Existuje zmäkčené pravidlo pre zjavne umelecké, satirické alebo fiktívne diela, kde zverejnenie musí prebehnúť len spôsobom, ktorý nekazí zážitok z diela: diskrétna poznámka v titulkoch namiesto pruhu cez celý film. A áno, aj tu sa objavuje výnimka pre presadzovanie práva.

Otvorene som hovoril, že sa mi tento dizajn nepáči. Deep fake sú každým rokom lacnejšie: menej referenčných dát, žiadna technická zručnosť, pár dolárov mesačne. Medzitým hlasové overovanie v bankovníctve a telefonické podvody „mami, potrebujem peniaze“ ukazujú presne, kam to smeruje. Štítok o zverejnení zmôže veľmi málo proti niekomu, koho celý biznis model je klamanie.

EÚ odvtedy tento bod čiastočne uznala. Digitálny omnibus z roku 2026 pridáva do článku 5 nový zákaz: AI generované **intímne zobrazenia bez súhlasu** („vyzliekacie“ aplikácie) a **materiál zobrazujúci sexuálne zneužívanie detí**, uplatniteľný po prechodnom období do **2. decembra 2026**. Takže jediná najhoršia kategória deep fake sa presúva zo „zverejnite to“ na „priamo zakázané“, so zodpovednosťou poskytovateľov tam, kde je takýto výstup odôvodnene predvídateľným a reprodukovateľným výsledkom ich systému, a nasadzujúcich subjektov, ktoré systémy úmyselne zneužijú. Všetko ostatné (falošný politik, falošná propagácia produktu) ostáva v režime zverejňovania. Či to vydrží, uvidíme.

Druhý pododsek článku 50 ods. 4 sa ľahko prehliadne a naozaj na ňom záleží. Nasadzujúce subjekty, ktoré používajú AI na generovanie alebo manipuláciu **textu uverejneného s cieľom informovať verejnosť o záležitostiach verejného záujmu**, to

musia zverejniť tiež. AI napísané správy sú pokryté, nielen AI vyrobené video. Ale existuje výnimka, ktorá v praxi urobí veľa práce: povinnosť sa neuplatňuje tam, kde AI vygenerovaný text prešiel ľudskou alebo redakčnou kontrolou a fyzická alebo právnická osoba zaň nesie redakčnú zodpovednosť. Zrozumiteľne: redakcia, kde redaktor AI návrh skontroluje, schváli a verejne sa zaň postaví, nedlhuje žiadny AI štítok. Web automaticky publikujúci AI napísané „správy“ bez človeka, ktorý by za ne zodpovedal, áno. Pravidlo je menej o tom, ako slová vznikli, a viac o tom, či sa za ne zodpovedá človek.

Drobné písmo, ktoré to spája

Článok 50 ods. 5 stanovuje spôsob a moment: všetky tieto zverejnenia musia byť jasné, rozlíšiteľné a doručené najneskôr pri *prvej* interakcii alebo vystavení, nie zahrabané v odseku 14 podmienok používania, a musia spĺňať požiadavky na prístupnosť. Článok 50 ods. 6 dodáva, že nič z toho nenahrádza kapitolu III: vysokorizikový systém, ktorý sa zároveň rozpráva so zákazníkmi, dlhuje obe sady povinností.

Keďže dátumy sú miesto, kde sa táto kapitola dá najľahšie pokaziť, tu je časová os na jednom mieste:

DÁTUM	ČO SA UPLATŇUJE
2. augusta 2026	Všetky štyri povinnosti podľa článku 50 (omnibus ich neodložil)
2. decembra 2026	Koniec odkladu označovania pre generatívne systémy, ktoré už boli na trhu k 2. aug. 2026
2. decembra 2026	Nový zákaz AI generovaných intímnych zobrazení bez súhlasu a materiálu zobrazujúceho zneužívanie detí podľa článku 5 sa začína uplatňovať
2. decembra 2027	Pre porovnanie: povinnosti pre vysokorizikové systémy z prílohy III. Toto sa posunulo, nie článok 50

Obvyklá výhrada: berte to s rezervou a konečné texty si overte so svojim právnym zástupcom. Omnibus je vo vestníku a usmernenia k článku 50 sú konečné, takže

dátumy vyššie sú tie, ktoré zaväzujú; ale usmernenia sa revidujú a členské štáty ešte len stavajú svoje orgány.

Overte si

1. Webový chatbot vašej banky odpovedá na otázky zákazníkov. Kto podľa článku 50 ods. 1 musí zabezpečiť, aby sa zákazníkom povedalo, že hovoria so systémom AI? A) Nasadzujúci subjekt chatbota B) Poskytovateľ, tým, že systém navrhne tak, aby používatelia boli informovaní C) Vnútroštátny orgán dohľadu nad trhom D) Nikto, lebo zverejnenie pri chatbotoch je do decembra 2027 dobrovoľné

2. Marketingová manažérka napíše vlastný text tlačovej správy a požiada generatívny AI nástroj len o opravu gramatiky a preklepov. Potrebuje výstup podľa článku 50 ods. 2 strojovo čitateľné označenie? A) Áno, čokoľvek, čo prejde cez systém AI, je syntetický obsah B) Áno, ale len ak text vyzerá realisticky C) Nie, lebo pomocné štandardné úpravy, ktoré podstatne nemenia vstup ani jeho význam, sú mimo povinnosti označovania D) Nie, lebo text článok 50 ods. 2 nikdy nepokrýva

3. Čo urobila novela digitálny omnibus z roku 2026 s povinnosťami v oblasti transparentnosti podľa článku 50? A) Odložila ich na december 2027 spolu s pravidlami pre vysokorizikové systémy B) Ponechala dátum 2. augusta 2026 a pridala len odklad do decembra 2026 pre označovanie generatívnymi systémami, ktoré už boli na trhu C) Podmienila ich uverejnením harmonizovaných noriem D) Zrušila povinnosť zverejňovať deep fake

4. Maloobchodný reťazec prevádzkuje softvér, ktorý z kamier v predajni odhaduje emócie zákazníkov (mimo akéhokoľvek zakázaného prostredia). Čo vyžaduje článok 50 ods. 3? A) Nič, pokiaľ sa záznamy do 24 hodín vymažú B) Posudzovanie zhody pred nasadením C) Poskytovateľ musí kamerový záznam opatriť vodoznakom D) Nasadzujúci subjekt musí vystavené osoby informovať, že systém je v prevádzke, a s osobnými údajmi zaobchádzať v súlade s GDPR

5. Spravodajský portál uverejní článok o výsledkoch volieb, ktorý navrhol systém AI a skontroloval ľudský redaktor, ktorý zaň nesie redakčnú zodpovednosť. Musí portál podľa článku 50 ods. 4 článok označiť ako

vygenerovaný AI? A) Nie, lebo výnimka pre ľudskú kontrolu a redakčnú zodpovednosť povinnosť zverejnenia odstraňuje B) Áno, všetok AI vygenerovaný text musí vždy niesť označenie C) Áno, politické témy nikdy nemôžu využiť redakčnú výnimku D) Nie, lebo text je pokrytý len vtedy, keď predstavuje deep fake

Odpovede: 1. B: článok 50 ods. 1 kladie povinnosť na poskytovateľov, ktorí musia systém navrhnuť a vyvinúť tak, aby fyzické osoby boli informované, pokiaľ to nie je zrejmé z kontextu. 2. C: článok 50 ods. 2 výslovne vylučuje AI vykonávajúcu pomocnú funkciu pri štandardných redakčných úpravách alebo podstatne nemeniacu vstupné údaje či ich sémantiku. 3. B: článok 50 nebol odložený a uplatňuje sa od 2. augusta 2026; omnibus poskytol len odklad označovania do 2. decembra 2026 pre už existujúce generatívne systémy. 4. D: článok 50 ods. 3 ukladá nasadzujúcim subjektom systémov na rozpoznávanie emócií a biometrickú kategorizáciu povinnosť informovať vystavené osoby a spracúvať osobné údaje podľa GDPR a súvisiacich pravidiel. 5. A: druhý pododsek článku 50 ods. 4 povinnosť zverejnenia neuplatňuje tam, kde text prešiel ľudskou alebo redakčnou kontrolou a niekto nesie redakčnú zodpovednosť za uverejnenie.

Povinnosť označovania, ktorú ste práve stretli, už pomenovala systémy, ktoré budú dominovať d'alšej kapitole: modely AI na všeobecné účely, ChatGPT tohto sveta, a osobitný režim, ktorý Akt postavil práve pre ne.

AI na všeobecné účely

Všetko v tejto knihe doteraz začínalo od tej istej otázky: čo tento systém *robí*? Triedič životopisov zoraďuje uchádzačov o prácu, takže pristane v prílohe III. Náladová kamera sleduje tvár recepčnej, takže naráža na pravidlá o rozpoznávaní emócií. Celý rebrík rizika predpokladá, že viete ukázať na účel a klasifikovať ho.

Teraz to skúste s ChatGPT. Ráno prekladá zmluvy, na obed píše rozchodovú správu, popoludní hrá personalistu sediaceho oproti vám na skúšobnom pracovnom pohovore a večer ponúkne priateľný názor na vašu kožnú vyrážku. Aký je jeho účel? Všetky. Žiadny konkrétny. Rebrík založený na účele, na ktorý sme liezli tri kapitoly, jednoducho nemá priečku pre systém, ktorý dokáže takmer čokoľvek.

Odpoveďou EÚ je kapitola V Aktu: samostatný režim mierený priamo na veľmi veľké modely za týmito nástrojmi, tie, ktorým odvetvie hovorí základné modely (foundation models) a Akt **modely AI na všeobecné účely**, skrátene GPAI. Táto kapitola je o tomto režime: kto doň spadá, čo dlhuje každý poskytovateľ v ňom a druhý rebrík vnútri neho, ktorý stúpa od bežného GPAI po GPAI so systémovým rizikom. Ten rebrík zrkadlí rebrík vysokého rizika, ktorý už poznáte, ale nikdy sa s ním nesmie zameniť, a keď tam prideme, budeme na to opatrní.

Model nie je systém

Pred všetkým ostatným jedno rozlíšenie, na ktorom stojí zvyšok tejto kapitoly. ChatGPT nie je model. ChatGPT je *aplikácia* (chatové rozhranie, bezpečnostné filtre, stránka s predplateným) postavená na modeli menom GPT. Model je surový motor; systém je auto postavené okolo neho. Akt tieto dve veci drží oddelene s takmer náboženskou disciplínou, a vy by ste mali tiež.

Tu je definícia z článku 3 bodu 63, a je to jedno z mála miest, kde si presné znenie zaslúži vlastný rámček.

„model AI na všeobecné účely“ je model AI, a to aj model AI trénovaný veľkým množstvom údajov s použitím samokontroly vo veľkom rozsahu, ktorý má významnú všeobecnú povahu a je schopný kompetentne vykonávať širokú škálu odlišných úloh bez ohľadu na spôsob, akým sa model uvádza na trh, a ktorý možno integrovať do rôznych nadväzujúcich systémov alebo aplikácií...

Zrozumiteľne: je to motor, ktorý (1) dokáže kompetentne robiť veľa rôznych vecí, nie jednu vec, a (2) je postavený na zabudovanie do produktov iných ľudí. Na oboch polovicách záleží. Model, ktorý predpovedá len nesplácanie úverov, nech je akokoľvek sofistikovaný, nie je na všeobecné účely; robí jednu prácu. A časť „integrovat' do nadväzujúcich systémov“ je tiché srdce definície: skutočný biznis týchto modelov nie je chatové okno, do ktorého píšete, ale tisíce aplikácií, ktoré na nich stavajú iné firmy. Startup pripájajúci GPT do produktu na medicínske triedenie pacientov je presne scenár, ktorý mali autori na mysli.

Definícia tiež spomína *samokontrolu vo veľkom rozsahu* (self-supervision), čo si zaslúži tridsať sekúnd. Pri učení s učiteľom bankové historické dáta modelu povedia, ktorí klienti svoje úvery splatili; označenia sú učiteľom. Veľké jazykové modely používajú iný trik. Vezmite akúkoľvek vetu z akejkoľvek knihy, povedzme „mačka sedela na...“, skryte posledné slovo a požiadajte model, aby ho uhádol. Pôvodný text odpoveď obsahuje, takže model sa vie oznámkovať sám, bez ľudského označovania. Urobte to naprieč slušnou časťou písaného internetu a dostanete plynulého univerzalistu, s ktorým sme začali. Dáta sú učiteľom samy sebe; odtiaľ samokontrola.

Ešte jedna definícia, keď už sme tu: článok 3 bod 66 definuje **systém AI na všeobecné účely** ako systém AI postavený na modeli AI na všeobecné účely. ChatGPT je systém; GPT je model. Kapitola V reguluje *model* a jeho poskytovateľa. Vo chvíli, keď niekto na ňom postaví systém (vrátane samotného tvorca modelu), na ten systém sa ako obvykle vzťahujú pravidlá pre systémy zo zvyšku Aktu.

Čo dlhuje každý poskytovateľ GPAI: článok 53

Predpokladajme, že váš model sa kvalifikuje ako na všeobecné účely. Čo dlhujete? Článok 53 vymenúva štyri povinnosti a tu je dátum, ktorý sa oplatí pripnúť: uplatňujú sa od **2. augusta 2025**. Novela digitálny omnibus z roku 2026, ktorá posunula termíny pre vysokorizikové systémy na december 2027 a august 2028, nechala túto kapitolu nedotknutú: tieto povinnosti zaväzujú dnes. (Čo sa 2. augusta 2026 zaplo, je formálna vrstva vymáhania: právomoc úradu pre AI vynútiť si informácie, nariadiť hodnotenia a uložiť pokuty špecifické pre GPAI až do 15 miliónov eur alebo 3 percent celosvetového obratu podľa článku 101. Dovtedy dohľad bežal ako zdvorilé, ale adresné „dialógy o súlade“, a k augustu 2026 je to stále tá istá tónina: výmeny Komisie s OpenAI a Anthropicom z leta 2026 o kyberútokoch spojených s ich modelmi boli zdieľaním informácií, nie formálnym konaním.)

Štyri povinnosti:

Technická dokumentácia pre regulátora. Vypracovať a udržiavať dokumentáciu modelu pokrývajúcu proces tréningu a testovania, výsledky hodnotenia a minimálny obsah uvedený v prílohe XI, pripravenú na odovzdanie úradu pre AI alebo vnútroštátnym orgánom na požiadanie (článok 53 ods. 1 písm. a)).

Dokumentácia pre ľudí, ktorí na vás stavajú. Pripraviť a udržiavať aktuálne informácie pre poskytovateľov systémov AI, ktorí váš model integrujú, aby dostatočne rozumeli jeho schopnostiam a obmedzeniam na splnenie vlastných povinností (článok 53 ods. 1 písm. b), minimálny obsah v prílohe XII). Ak si pamätáte kapitolu o vysokom riziku, toto sa rýmuje: dokumentácia dovnútra plus návod von. Ale všimnite si, ako hlboko rým siaha. Vo svete systémov nabiehajú dokumentačné povinnosti až pri vysokom riziku. Vo svete modelov ich nesie *každý* model na všeobecné účely, so systémovým rizikom či bez. Latka tu sedí o úroveň nižšie.

Politika autorského práva. Zaviesť politiku súladu s autorským právom EÚ vrátane rešpektovania strojovo čitateľných výhrad nositeľov práv voči vyťažovaniu textov a dát (článok 53 ods. 1 písm. c)). Kapitola 12 sa do príbehu autorského práva ponorí poriadne; nateraz vedzte, že povinnosť existuje.

Verejné zhrnutie tréningových dát. Zverejniť „dostatočne podrobné zhrnutie“ obsahu použitého na tréning modelu podľa šablóny od úradu pre AI (článok 53 ods. 1 písm. d)). Keď som to čítal prvýkrát, zdvihol som obočie: veľké laboratória si recepty

na tréningové dáta žiarlivo strážia a tu je zákonná povinnosť verejne ich zhrnúť. Šablóna už nie je hypotetická: úrad pre AI zverejnil povinnú verziu s vysvetľujúcim oznámením v decembri 2025. Vyžaduje okrem iného zverejnenie najväčších webových domén v tréningových dátach (horných 10 percent domén podľa objemu; menšie firmy dostávajú ľahšiu verziu 5 percent alebo 1 000 domén). Jedno prechodné ustanovenie ranu zmierňuje: modely, ktoré už boli na trhu pred 2. augustom 2025, majú na dobehnutie tejto dokumentácie čas do 2. augusta 2027. Nové modely sú v súlade od prvého dňa.

Nad rámec tých štyroch článok 53 pridáva očakávanú údržbu: spolupracovať s Komisiou a vnútroštátnymi orgánmi, a všetko, čo odovzdáte (vrátane obchodných tajomstiev), je chránené pravidlami Aktu o dôvernosti.

Open-source zľava

Článok 53 ods. 2 vyrezáva skutočný ústupok. Ak svoj model vydáte pod bezplatnou licenciou s otvoreným zdrojovým kódom a skutočne otvorene (parametre, váhy, architektúra a informácie o používaní všetky verejne dostupné), potom sa dve dokumentačné povinnosti, písm. a) a b), na vás nevzťahujú. Logika je dosť férová: transparentnosť, ktorú by tie dokumenty poskytli, už sedí na internete pre kohokoľvek na preskúmanie.

Všimnite si, čo zľava *nepokrýva*. Politika autorského práva a verejné zhrnutie tréningových dát stále platia, lebo otvorenosť o váhach nie je otvorenosť o tom, na čom ste trénovali. A celá výnimka sa vyparí vo chvíli, keď model nesie systémové riziko. Open-source vydanie modelu na špici nevykúpi jeho poskytovateľa z ťažkej úrovne. Vzhľadom na to, že niektoré veľké laboratória vydávajú svoje najschopnejšie modely otvorene, tá výnimka nie je poznámka pod čiarou; je to klauzula, ktorá robí skutočnú prácu.

Klasifikácia ako systémové riziko: predpoklad, nie verdikt

Kapitola V má vlastnú vnútornú eskaláciu a jej tvar bude pôsobiť známo. Tak ako systém AI môže vyliezť z „bežného“ na „vysokorizikový“, model GPAI môže vyliezť z „bežného GPAI“ na **GPAI so systémovým rizikom**: úroveň vyhradenú pre skutočné

modely na špici, s pripojenými ťažšími povinnosťami. Tvar je známy, ale rebríky nie sú ten istý rebrík, a tu chcem byť pevný, lebo práve tu vidím ľudí kĺzať.

POZOR

Vysoké riziko je vlastnosť *systémov AI*: žije v kapitole III a závisí od toho, na čo sa systém používa (nábor, úvery, kontrola hraníc). **Systémové riziko** je vlastnosť *modelov*: žije v kapitole V a závisí od toho, aký schopný je surový motor, bez ohľadu na použitie. Model na špici so systémovým rizikom nie je „vysokorizikový“ a nazývať ho tak nie je neškodná skratka; ukazuje vás na nesprávnu sadu článkov. Systémy môžu byť vysokorizikové; modely môžu niesť systémové riziko.

Ako teda model lezie po tomto druhom rebríku? Článok 51 dáva dve cesty.

Prvá je založená na schopnostiach: model má „spôsobilosti s veľkým vplyvom“, posudzované technickými nástrojmi a benchmarkmi. A keďže „veľký vplyv“ sa ťažko meria priamo, článok 51 ods. 2 dodáva zástupnú mieru: model sa **predpokladá** za majúci spôsobilosti s veľkým vplyvom, keď jeho tréningovanie použilo viac než 10^{25} operácií s pohyblivou rádovou čiarkou výpočtového výkonu. Nemusíte to číslo precítiť; len vedzte, že je to obrovské množstvo výpočtov, také, aké vierohodne dosiahla len hrstka tréningových behov na špici. Na kalibráciu: usmernenia Komisie ku GPAI berú zhruba 10^{23} FLOP ako orientačné meradlo pre to, aby model *vôbec bol* na všeobecné účely. Čiara systémového rizika sedí stokrát vyššie. Väčšina komerčných modelov vrátane mnohých známych žije pohodlne pod ňou; táto úroveň bola postavená pre špičku, nie pre hlavný prúd.

Slovo, ktoré v tej vete robí právnu prácu, je *predpokladá*. Prekročenie 10^{25} neopečiatkuje váš model „systémové riziko“; presúva bremeno na vás. Podľa článku 52 ods. 2 môže poskytovateľ predložiť podložené argumenty, že napriek výpočtovému výkonu model v skutočnosti systémové riziká nepredstavuje, a Komisia ich prijme alebo zamietne (článok 52 ods. 3). Vyvrátiteľný predpoklad, právnickou rečou. Toto je jediný najviac skreslene hlásený fakt o kapitole V, tak jasne: 10^{25} FLOP je miesto, kde sa rozhovor začína, nie kde končí.

Druhá cesta je určenie: Komisia môže model klasifikovať ako so systémovým rizikom z vlastnej iniciatívy alebo po upozornení svojho vedeckého panelu, podľa kritérií v prílohe XIII: počet parametrov, veľkosť súboru dát, tréningový výpočtový výkon, vstupné a výstupné modalities, výsledky benchmarkov, dokonca počet registrovaných používateľov (článok 51 ods. 1 písm. b), článok 52 ods. 4). Takže aj model, ktorý nikdy neprekročil výpočtovú čiaru, môže byť vytiahnutý po rebríku, ak to jeho skutočné schopnosti a dosah odôvodňujú. Určený poskytovateľ môže požiadať o prehodnotenie, najskôr šesť mesiacov po rozhodnutí (článok 52 ods. 5).

Ktoré modely boli naozaj klasifikované? Tu je úprimná odpoveď a možno vás prekvapí: **nikto mimo procesu to nevie s istotou.** Článok 52 ods. 6 ukladá Komisii zverejniť zoznam modelov so systémovým rizikom, ale k augustu 2026 žiadny takýto verejný register neexistuje. Odhady o tom, ktoré modely na špici prekračujú 10^{25} FLOP, kolujú neustále, ale tréningový výpočtový výkon sa zriedka zverejňuje, takže sú to presne to: odhady. Nebudem menovať a s každým, kto sebavedome menuje, by som zaobchádzal podozrieľavo.

MÔJ POHĽAD

Čo by mal poskytovateľ v mierke špičky predpokladať, je konzervatívne čítanie: ak je váš tréningový beh kdekoľvek blízko čiary, počítajte s tým, že sa kvalifikujete. Radšej by som, aspoň nateraz, počítal s horšou možnosťou.

Dva týždne na zdvihnutie ruky

Je tu jedna procesná povinnosť so skutočnými zubami, ľahko prehliadnuteľná, lebo sa skrýva v článku 52 ods. 1: poskytovateľ, ktorého model spĺňa výpočtovú podmienku, musí **oznámiť Komisii do dvoch týždňov**. Dva týždne od chvíle, keď je prah splnený, alebo od chvíle, keď sa stane známym, že splnený *bude*. Keďže tréningové behy sa plánujú a rozpočtujú dlho dopredu, tá druhá klauzula znamená, že hodiny môžu začať bežať skôr, než sa tréningovanie vôbec skončí. Oznámenie môže niesť vyvracajúce argumenty, o ktorých sme práve hovorili. Ak poskytovateľ mlčí, Komisia môže model jednoducho určiť aj tak, keď sa to dozvie. Táto povinnosť, ako zvyšok kapitoly, je živá od augusta 2025.

Ďalšie povinnosti na vrchole: článok 55

Model klasifikovaný so systémovým rizikom si ponecháva všetky povinnosti podľa článku 53 a pridáva štyri ďalšie (článok 55 ods. 1):

- **Hodnotenie modelu s adverzariálnym testovaním.** Vyhodnotiť model podľa najmodernejších protokolov vrátane štruktúrovaných pokusov donútiť ho správať sa zle (red-teaming, v reči odvetvia). Ak vás príbeh o otrávených tréningových dátach z kapitoly 1 naučil, že model môže niesť skryté režimy zlyhania, ktoré jeho tvorca nikdy nezamýšľal, toto je povinnosť, ktorá hovorí: choďte ich pred vydaním loviť a lov zdokumentujte.
- **Posúdiť a zmierniť systémové riziká na úrovni Únie:** riziká plynúce z vývoja, vydania alebo používania modelu, vystopované k ich zdrojom.
- **Hlásenie závažných incidentov.** Sledovať, dokumentovať a bez zbytočného odkladu hlásiť závažné incidenty a prijaté nápravné opatrenia úradu pre AI a vnútroštátnym orgánom. Aj to dostalo praktické lešenie: Komisia zverejnila šablónu na hlásenie incidentov v novembri 2025.
- **Kybernetická bezpečnosť**, modelu aj fyzickej infraštruktúry, na ktorej beží. Váhy modelov na špici patria k najcennejším súborom na zemi; Akt trvá na tom, aby sa podľa toho bránili.

Všimnite si tóninu, v akej sú tieto povinnosti napísané: hodnotiť, zmierňovať, hlásiť, zabezpečiť. Nie „nevydávajte“: kapitola V nikdy žiadny model nezakazuje. Predpokladá, že modely na špici budú existovať, a od ich tvorcov žiada, aby ich spravovali ako kritickú infraštruktúru, ktorou sa stávajú.

Kódex postupov a kto ho podpísal

Ako poskytovateľ toto všetko vlastne preukáže? Článok 53 ods. 4 a článok 55 ods. 2 načrtávajú hierarchiu dôkazov a mať ju správne záleží viac, než sa zdá.

Na vrchole sedia **harmonizované normy**: dodržte jednu a získate formálny *predpoklad zhody* pre to, čo norma pokrýva. O úroveň nižšie sedia **kódexy postupov** podľa článku 56: ich dodržiavanie vám umožní *preukázať súlad*, kým prídu harmonizované normy. To je skutočne užitočný právny štít, ale slabší. Kódex ukazuje,

že ste v súlade; norma vedie regulátora k tomu, že to predpokladá. A poskytovateľ, ktorý nedodrížiava ani jedno, musí Komisii priamo dokázať „alternatívne primerané prostriedky súladu“, čo je nepohodlné kreslo.

STAV K AUGUSTU 2026

V Úradnom vestníku je citovaných **nula** harmonizovaných noriem podľa AI Actu, takže cesta predpokladu zhody existuje len na papieri a kódex postupov je jediná hra v meste. (A nie, certifikát ISO/IEC 42001 ho nenahrádza: užitočný pre váš systém riadenia, mlčiaci o zhode s AI Actom.) Aj zoznamy signatárov sa hýbu; Komisia udržiava aktuálny na svojom webe o digitálnej stratégii. Kým sa na tento snímok spoľahnete, znovu overte oboje.

Ten kódex existuje a jeho krátka história je lepšou lekciou o tom, ako regulácia technológií v EÚ naozaj funguje, než akýkoľvek text článku. **Kódex postupov pre GPAI** bol zverejnený 10. júla 2025 v troch kapitolách (Transparentnosť, Autorské právo, Bezpečnosť a ochrana) plus formulár dokumentácie modelu a Komisia ho označila za primeraný približne v čase, keď v auguste 2025 povinnosti ožili. Potom prišla tá zaujímavá časť: kto podpíše dobrovoľný nástroj? Zhruba dve desiatky poskytovateľov áno, vrátane väčšiny najväčších mien: Amazon, Anthropic, Google, IBM, Microsoft, Mistral, OpenAI. **Meta verejne odmietla** s argumentom, že kódex ide nad rámec samotného Aktu. **xAI podpísalo presne jednu kapitolu**, Bezpečnosť a ochrana, a ostatné dve nechalo na stole, možnosť à la carte, ktorú kódex pre svoju bezpečnostnú kapitolu pripúšťa. Veľké čínske laboratória sa väčšinou držali bokom.

Vzorec stojí za zapamätanie: aj „dobrovoľný“ kódex sa stáva strategickým rozhodnutím, keď je alternatívou dokazovať súlad regulátorovi tou ťažkou cestou. Meta odmietnutím podpisu článku 53 neunikla; len si vybrala nepohodlné kreslo.

Overte si

1. Vaša firma postaví chatbota na medicínske rady na veľkom modeli licencovanom od veľkého AI laboratória. Ktorý štítok podľa kapitoly V sedí na čo? A) Váš chatbot je model GPAI; model laboratória je systém GPAI B) Model

laboratória je model GPAI; váš chatbot je systém AI postavený na ňom C) Oba sú modely GPAI so systémovým rizikom D) Ani jeden nie je pokrytý AI Actom do roku 2027

2. Ktorá povinnosť podľa článku 53 sa stále vzťahuje na skutočne bezplatný open-source model GPAI (váhy, architektúra a informácie o používaní všetky verejné), za predpokladu, že nemá systémové riziko? A) Technická dokumentácia pre úrad pre AI podľa článku 53 ods. 1 písm. a) B) Dokumentácia pre nadväzujúcich poskytovateľov systémov podľa článku 53 ods. 1 písm. b) C) Verejné zhrnutie tréningových dát podľa článku 53 ods. 1 písm. d) D) Žiadna, lebo open-source modely sú z kapitoly V úplne vyňaté

3. Nový tréningový beh poskytovateľa prekročí 10^{25} FLOP. Aký je právny dôsledok? A) Model je automaticky a s konečnou platnosťou klasifikovaný ako so systémovým rizikom a objaví sa vo verejnom registri B) Model sa predpokladá za majúci spôsobilosti s veľkým vplyvom; poskytovateľ musí do dvoch týždňov oznámiť Komisii a môže predložiť argumenty vyvracajúce klasifikáciu C) Nič, pokiaľ Komisia model neurčí aj podľa prílohy XIII D) Model sa stáva vysokorizikovým podľa kapitoly III

4. Ktoré tvrdenie o kódexe postupov pre GPAI je k polovici roka 2026 správne? A) Jeho dodržiavanie dáva formálny predpoklad zhody B) Nikdy nebol dokončený, takže Komisia uložila spoločné pravidlá vykonávacím aktom C) Umožňuje signatárom preukázať súlad, kým nebudú uverejnené harmonizované normy, a podpísali ho zhruba dve desiatky poskytovateľov, hoci nie všetci D) Jeho podpísanie je povinné pre každého poskytovateľa GPAI pôsobiaceho v EÚ

5. Ktorá povinnosť patrí poskytovateľom modelov GPAI so systémovým rizikom nad rámec bežných povinností podľa článku 53? A) Registrácia modelu v databáze EÚ pre vysokorizikové systémy B) Vykonanie posudzovania zhody s notifikovanou osobou C) Adverzariálne testovanie, zmiernovanie systémových rizík, hlásenie závažných incidentov a kybernetická bezpečnosť podľa článku 55 D) Vymenovanie ľudského dozorca pre každú nadväzujúcu aplikáciu

Odpovede. 1: B. Motory typu GPT sú *modely* GPAI; čokoľvek sa na nich postaví (laboratóriom alebo vami) je *systém*, ktorý sa riadi zvyškom Aktu. 2: C. Open-source

výnimka v článku 53 ods. 2 sníma len dve dokumentačné povinnosti a) a b); politika autorského práva a verejné zhrnutie tréningových dát ostávajú a celá výnimka mizne pri modeloch so systémovým rizikom. 3: B. Článok 51 ods. 2 vytvára vyvrátiteľný predpoklad a články 52 ods. 1 a 52 ods. 2 stanovujú dvojťždňové oznámenie s možnosťou argumentovať proti klasifikácii; k augustu 2026 žiadny verejný register určených modelov neexistuje. 4: C. Kódexy postupov preukazujú súlad (články 53 ods. 4 a 55 ods. 2); len harmonizované normy dávajú predpoklad zhody a v Úradnom vestníku zatiaľ nebola citovaná žiadna; Meta odmietla podpísať a xAI podpísalo len bezpečnostnú kapitolu. 5: C. Článok 55 pridáva hodnotenie s adverzariálnym testovaním, zmierňovanie rizík na úrovni Únie, hlásenie incidentov a kybernetickú bezpečnosť; registrácia v databáze a posudzovanie zhody notifikovanou osobou patria do režimu vysokorizikových *systémov*, iného rebríka.

Ďalej: kto to všetko naozaj vymáha, od úradu pre AI a vnútroštátnych orgánov po pokuty a úprimný stav časovej osi.

Správa a riadenie, sankcie, časová os

Desať kapitol vám niekto hovorí, čo *musíte* robiť. Zdokumentovať úverový model svojej banky. Zverejniť svojho chatbota. Zložiť náladovú kameru z recepčnej. Celý ten čas sa pravdepodobne hromadí férová otázka: kto to vlastne hovorí? Ak to vaša firma pokazí, kto zaklope na dvere: bruselský inšpektor, váš národný telekomunikačný úrad, nejaká nová agentúra pre AI? A dostal už niekto, niekde, naozaj pokutu?

Táto kapitola na tieto otázky odpovedá úprimne, čo znamená, že odpovede budú menej dramatické než titulky. Začneme podpornou rukou Aktu, regulačnými experimentálnymi prostrediami (sandboxmi): sú široko nepochopené, tak ich vezmeme pomaly a poviem vám otvorene, čo si o nich myslím. Potom stretneme inštitúcie, ktoré Akt naozaj prevádzkujú, pozrieme sa na úrovne pokút a porovnáme zákon na papieri s realitou vymáhania v polovici roka 2026. Uzavrieme jedinou tabuľkou, ktorú sa oplatí vytlačiť, plným kalendárom uplatňovania.

Experimentálne prostredia: podpora, nie brána

Predstavte si fintech startup, dvanásť ľudí, staviajúci model úverovej bonity, jednoznačne vysokorizikový podľa prílohy III. Prečítali si kapitolu o požiadavkách a sú nervózni: riadenie rizík, správa dát, technická dokumentácia, všetko pre nich nové. Čo by milovali, je niekto zo strany regulátora, kto by im pozrel cez plece *skôr*, než budú stávky skutočné, a povedal „áno, to je zhruba to, čo sme mysleli“.

Na to je regulačné experimentálne prostredie pre AI. Koncept je požičaný z fintechu, kde regulátori roky prevádzkujú dozorované testovacie prostredia. Článok 57 ods. 5 opisuje experimentálne prostredie ako kontrolované prostredie, ktoré podporuje inovácie a umožňuje vám na obmedzený čas vyvíjať, trénovať, testovať a validovať inovatívny systém AI podľa konkrétneho plánu experimentálneho prostredia dohodnutého medzi vami a príslušným orgánom (celé obsadenie orgánov stretneme o chvíľu). Dostanete usmernenie k regulačným očakávaniam, dohľad, pomoc s identifikáciou rizík. MSP a startupy získavajú prístup bezplatne (článok 58) a orgány musia o vašej žiadosti rozhodnúť do troch mesiacov.

Teraz časť, ktorú ľudia chápu zle, tak budem priamy.

POZOR

Experimentálne prostredie nie je schvaľovacia brána. Na uvedenie vysokorizikového systému na trh nepotrebuje jeho povolenie: prístup vedie cez posudzovanie zhody podľa článku 43, presne ako je opísané skôr v tejto knihe, či ste sa k experimentálnemu prostrediu kedy priblížili, alebo nie. Účasť je dobrovoľná podpora inovácií, mierená najmä na menších hráčov bez oddelenia regulačných záležitostí.

Čo z toho teda máte? Článok 57 ods. 7 to vypisuje:

AKT HOVORÍ · ČLÁNOK 57 ODS. 7

„...Orgány dohľadu nad trhom a notifikované osoby v tejto súvislosti pozitívne zohľadnia výstupné správy a písomný dôkaz, ktoré poskytol vnútroštátny príslušný orgán, aby sa v primeranom rozsahu urýchlili postupy posudzovania zhody.“

Zrozumiteľne: experimentálne prostredie vám dá písomný záznam, že príslušný orgán vás sledoval pri práci a videl, že pravidlá beriete vážne, záznam, ktorý urýchli vaše skutočné posudzovanie zhody, ale nenahradí ho. Ešte jeden skutočne cenný bonus sa skrýva v článku 57 ods. 12: dodržiavajte dohodnutý plán experimentálneho prostredia a usmernenia orgánu v dobrej viere a za porušenia AI Actu počas experimentovania sa neukladajú žiadne správne pokuty. Bezpečný priestor v doslovnom zmysle, hoci za škodu spôsobenú tretím stranám ostávate zodpovední podľa bežného práva.

Dve spresnenia, lebo tieto body sa mýlia aj v odborných kruhoch. Po prvé, **databáza EÚ pre vysokorizikové systémy nie je práca experimentálnych prostredí**. Tú databázu spravuje Komisia (článok 71) a poskytovatelia v nej svoje systémy registrujú priamo, pred uvedením na trh (článok 49). Experimentálne prostredia v tom nehrajú žiadnu rolu. Po druhé, **testovanie v reálnych podmienkach mimo experimentálneho prostredia**, teda testovanie vášho systému z prílohy III so skutočnými ľuďmi pred posudzovaním zhody, je samostatná koľaj podľa článku 60,

schvaľovaná orgánom dohľadu nad trhom: predložíte plán testovania a mlčanie 30 dní sa počíta ako tichý súhlas. Testovanie v reálnych podmienkach sa môže diať aj *vnútri* experimentálneho prostredia, tam pod dohľadom, ale potom podmienky sedia vo vašom pláne experimentálneho prostredia a žiadne 30-dňové hodiny nebežia.

Kde experimentálne prostredia stoja dnes? V sklze. Pôvodný termín, do ktorého mal mať každý členský štát aspoň jedno funkčné experimentálne prostredie, bol 2. august 2026; novela omnibus z roku 2026 ho posunula na 2. august 2027 a pridala aj experimentálne prostredie na úrovni EÚ prevádzkované úradom pre AI a rozšírila testovanie v reálnych podmienkach na systémy z prílohy I. Ak experimentálne prostredie vašej krajiny ešte neexistuje, nie je to problém súladu, lebo účasť nikdy nebola povinná. Ale podporná ruka Aktu prichádza neskôr než ruka povinností. Plánujte podľa toho.

Moje myšlienky o experimentálnych prostrediach

MÔJ POHĽAD

Priznám sa, že sklz som videl prichádzať. Bazén ľudí, ktorí rozumejú strojovému učeniu aj právu EÚ o výrobkoch, je plytký a dozorné experimentálne prostredie nemôžete obsadiť juniormi. Žiadať 27 krajín, aby takých ľudí najali na dvojročnej časovej osi, bolo vždy optimistické.

Ďalšie dve veci ma znepokojujú, keď sa experimentálne prostredia naozaj otvoria.

Prvá je pracovná záťaž. Videli ste, aká široká je definícia systému AI a ako široko zametajú kategórie prílohy III. Banky a poisťovne, ktoré školím, si už mapujú prípady použitia; v deň, keď sa v ich krajine otvorí dobre vedené experimentálne prostredie, naň udrie prvá vlna žiadostí a päťčlenná kancelária čeliaca tej vlne znamená rady. Trojmesačná lehota na rozhodnutie pomáha, len ak má kancelária za ňou kapacitu. A na oboch stranách je krivka učenia: firmy sa musia naučiť, čo orgán očakáva, a orgán sa musí naučiť uplatňovať skutočne nový zákon na skutočne rôznorodé systémy. Ani jedna strana sa nemá o čo oprieť v desaťročiach ustálenej praxe.

Druhá je kultúra. Experimentálne prostredie a testovanie v reálnych podmienkach vôbec predpokladajú, že pracujete ako vedec: sformulujete hypotézu, navrhnete experiment, testujete v obmedzenom rozsahu, meriate. Medicína funguje takto, lebo nový liek nemá historické dáta a kontrolovaná štúdia je jediný poctivý dôkaz. Väčšina európskych firiem podľa mojej skúsenosti nie. Stavajú z historických dát a pozorovania; štruktúrované experimentovanie je návyk, ktorý si nikdy nevytvorili. Mnohé zistia, že ťažkou časťou experimentálneho prostredia nie je papierovanie, ale disciplína, ktorú predpokladá.

Berte to všetko ako čítanie jedného konzultanta, nie proroctvo. Úprimné zhrnutie: experimentálne prostredia sú dobrá myšlienka prichádzajúca pomaly a najviac z nich vyťažia firmy, ktoré ich budú brať ako kanál učenia, nie ako rad, v ktorom sa stojí.

Kto to vlastne prevádzkuje

Experimentálne prostredie vás podporuje; nepolicajtuje vás. Kto teda áno? AI Act je nariadenie EÚ, ale EÚ nepostavila jedinú AI políciu. Vymáhanie je rozdelené medzi dve úrovne a vedieť, ktorá rieši čo, vás ušetrí od e-mailovania nesprávnej inštitúcie.

Na národnej úrovni: vaše príslušné orgány. Článok 70 vyžaduje, aby každý členský štát určil aspoň jeden *notifikujúci orgán* (ktorý akredituje notifikované osoby z kapitoly o posudzovaní zhody) a aspoň jeden *orgán dohľadu nad trhom* (ktorý kontroluje systémy na trhu a vyšetruje, keď niečo vyzerá zle). Orgán dohľadu nad trhom je ten, na ktorom záleží pre vaše každodenné riziko. Ak konkurent nahlási váš dashboard rozpoznávania emócií alebo novinár napíše o vašom bodovacom systéme, toto je orgán, ktorý otvorí spis, vrátane zákazov podľa článku 5. Každá krajina určí jeden z týchto orgánov ako svoje jednotné kontaktné miesto.

Na úrovni EÚ: úrad pre AI. Úrad pre AI (článok 64), umiestnený vnútri Európskej komisie, je najbližšie k centrálnemu mozgu, aký Akt má. Priamo dozerá na modely AI na všeobecné účely; povinnosti GPAI z predchádzajúcej kapitoly sú jeho teritórium, nie územie vášho národného regulátora. A koordinuje všetko ostatné: usmernenia, kódexy postupov, podporu vnútroštátnych orgánov.

Okolo nich tri podporné orgány. Európska rada pre umelú inteligenciu (články 65 a 66), s jedným zástupcom za každý členský štát, existuje preto, aby sa 27 národných

praxí vymáhania nerozchádzalo. Vedecký panel nezávislých expertov (článok 68) radí v otázkach GPAI vrátane upozorňovania úradu pre AI na systémové riziká; bol ustanovený v roku 2025 so šesťdesiatimi expertmi. A poradné fórum (článok 67) dáva formálny hlas priemyslu, MSP, akadémii a občianskej spoločnosti. Priamo s ktorýmkoľvek z nich budete rokovať zriedka, ale keď Komisia zverejní usmernenia, táto mašinéria je miesto, kde sa o nich hádalo.

A jedna vec, ktorú by ste mali naozaj ísť vyskúšať. Od októbra 2025 Komisia prevádzkuje **AI Act Service Desk** s jednotnou informačnou platformou: kontrolór súladu, interaktívny prehliadač Aktu a kanál na kladenie otázok. Je na ai-act-service-desk.ec.europa.eu, je zadarmo a odpovedá na otázku „koho sa vôbec spýtať?“, ktorá raným rokom Aktu chýbala. Ak po tejto kapitole urobíte jednu akciu, strávte tam pätnásť minút s jedným z vlastných prípadov použitia. Odpovede sú nezáväzné, ale sú najbližšie k oficiálnemu stanovisku, aké dostanete bez najatia advokátskej kancelárie.

Cenník

Teraz čísla, ktoré každý cituje na konferenciách. Článok 99 stanovuje tri úrovne správnych pokút a v každom prípade je strop pevná suma *alebo* percento celkového celosvetového ročného obratu, podľa toho, čo je **vyššie**:

- **35 miliónov eur alebo 7 percent** za porušenie zákazov podľa článku 5. Nasadíte tú náladovú kameru na recepcnú a ste v tejto úrovni.
- **15 miliónov eur alebo 3 percentá** za porušenie hlavných prevádzkových povinností: povinností poskytovateľa podľa článku 16, povinností nasadzujúceho subjektu podľa článku 26, povinností dovozcov a distribútorov, pravidiel pre notifikované osoby a povinností v oblasti transparentnosti podľa článku 50. Tu by pristála väčšina realistických firemných zlyhaní; myslíte na banku, ktorá prevádzkuje svoj vysokorizikový úverový model bez požadovanej dokumentácie a dohľadu.
- **7,5 milióna eur alebo 1 percento** za poskytnutie nesprávnych, neúplných alebo zavádzajúcich informácií orgánom alebo notifikovaným osobám.

Tri zmierňovače stoja za poznanie. Pri MSP a startupoch sa každý strop preklopí na to, čo je *nížšie* zo sumy a percenta (článok 99 ods. 6), skutočná milosť pre dvanásťčlenný

fintech. Pokuty musia byť účinné, primerané a odrádzajúce a článok 99 ods. 7 vymenúva okolnosti, ktoré formujú skutočnú výšku: závažnosť a trvanie, spoluprácu, sebaoznámenie, úmysel verzus nedbanlivosť. A sú to *stropy*, nie cenovky; nič v Akte nehovorí, že prvé porušenie ide kamkoľvek blízko maxima.

Vedľa sedia dva osobitné režimy. Keď pravidlá porušia samotné inštitúcie EÚ, pokutuje ich európsky dozorný úradník pre ochranu údajov, s oveľa nižšími stropmi (do 1,5 milióna eur za zakázané praktiky, článok 100). A pre poskytovateľov modelov AI na všeobecné účely prichádzajú pokuty od samotnej Komisie, nie od vnútroštátnych orgánov: do 15 miliónov eur alebo 3 percent podľa článku 101. Jedna časová jemnosť: kapitola o sankciách sa uplatňuje od 2. augusta 2025, ale formálna právomoc Komisie pokutovať poskytovateľov GPAI podľa článku 101 zahryzne až od 2. augusta 2026.

Kto bežné pokuty naozaj ukladá? Členské štáty. Článok 99 ods. 1 vyžaduje, aby každá krajina stanovila vlastné pravidlá sankcií v rámci stropov Aktu, a tento detail vedie priamo k najmenej diskutovanému faktu o AI Acte v roku 2026.

Výsledková tabuľa, úprimne

Tu by som vám radšej dal nepríjemnú pravdu než pôsobivý slajd.

STAV K AUGUSTU 2026

So zákazmi vymáhateľnými od februára 2025 a plnou kapitolou o sankciách živou od augusta 2025 neexistuje nikde v Európskej únii potvrdená pokuta podľa AI Actu. Ani jedna. Ani potvrdené formálne konanie. Kým to zopakujete v zasadačke, znovu overte; môže sa to zmeniť jediným rozhodnutím.

Prečo? Pozrite sa na inštalatérčinu. Členské štáty mali určiť svoje príslušné orgány do 2. augusta 2025; do jari 2026 tak naozaj urobilo len približne 8 z 27. Pokuta potrebuje orgán, ktorý ju uloží, a národný zákon o sankciách, o ktorý sa oprie, a väčšina krajín nemá plne zavedené ani jedno. Taliansko tam bolo prvé a doteraz najkompletnejšie: zákon 132/2025, účinný od októbra 2025, menuje jeho orgány a je jediným úplným národným sadzobníkom sankcií v účinnosti. Slovensko doteraz prijalo len zákon č. 318/2025 Z. z., úzku novelu zákona o posudzovaní zhody účinnú od januára 2026; plnší

slovenský návrh, ktorý by z MIRRI urobil orgán dohľadu nad trhom a jednotné kontaktné miesto, prešiel 14. augusta 2026 prvým vládny kolom a mieri do Legislatívnej rady, potom na vládu a do parlamentu. Česko je tiež vo fáze návrhu: jeho adaptačný zákon by odovzdal dohľad nad trhom ČTÚ, ale neprešiel. Ak pracujete v Bratislave alebo Prahe, váš regulátor AI sa, celkom doslova, ešte len formuje.

Najbližšie k testovaciemu prípadu má Maďarsko, kde organizácie občianskej spoločnosti tvrdia, že rozšírené sledovanie rozpoznávaním tváří porušuje zákaz diaľkovej biometrickej identifikácie v reálnom čase podľa článku 5, a tlačili na Komisiu, aby konala, ale žiadne konanie o porušení nebolo potvrdené. A pokuty, ktoré ste možno videli v správach v spojení s AI firmami (pokuta 5 miliónov eur od talianskeho úradu na ochranu údajov pre tvorca chatbota Replika, nahromadené pokuty proti Clearview AI), sú pokuty podľa **GDPR**. Iný zákon, iné orgány. Držte režimy oddelene, lebo vaši kolegovia nebudú.

Takže sa môžete uvoľniť? Ja by som sa neuvolnil. Každý režim vymáhania v histórii EÚ, GDPR nevynímajúc, začínal presne touto tichou fázou a tichá fáza GDPR skončila pokutami v stovkách miliónov. Orgány, ktoré sa teraz určujú, zdedia backlog sťažností; maďarský spis ukazuje, že občianska spoločnosť ich už generuje. Moje úprimné čítanie: riziko vymáhania je pred vami, nie za vami. To nie je dôvod Akt ignorovať; je to vzácne, zatvárajúce sa okno, v ktorom môžete veci opraviť pokojne, namiesto pod vyšetrovaním.

Kalendár

Ešte jedna pohyblivá súčiastka pred tabuľkou: časovú os v roku 2026 pretvarovala novela digitálny omnibus, odsúhlasená Európskym parlamentom 16. júna a Radou 29. júna 2026. V Úradnom vestníku vyšla 24. júla 2026 ako nariadenie (EÚ) 2026/1744 a účinná je od 27. júla 2026. Jej titulným účinkom bolo odloženie termínov pre vysokorizikové systémy zhruba o rok či viac. Jednu vec **neurobila**, napriek pretrvávajúcim fámam: neodložila povinnosti v oblasti transparentnosti podľa článku 50, ktoré sa uplatňujú od 2. augusta 2026, ako bolo pôvodne naplánované. A nové dátumy sú pevné kalendárne dátumy. Skorší návrh podmieniť ich pripravenosťou harmonizovaných noriem bol zamietnutý, tak nedovoľte nikomu, aby vám hovoril, že termíny „začnú až keď prídu normy“. (V Úradnom vestníku je stále citovaných nula

harmonizovaných noriem a nie, certifikát ISO/IEC 42001 nie je súlad s AI Actom; ale hodiny bežia bez ohľadu na to.)

Tu je plný obraz, ako stojí v auguste 2026:

DÁTUM	ČO SA UPLATŇUJE
1. augusta 2024	AI Act nadobúda účinnosť (nariadenie (EÚ) 2024/1689)
2. februára 2025	Zakázané praktiky (článok 5); povinnosť gramotnosti v oblasti AI (článok 4)
2. augusta 2025	Povinnosti pre modely GPAI; orgány správy a riadenia; sankcie (kapitola XII); termín na určenie vnútroštátnych orgánov
2. augusta 2026	Povinnosti v oblasti transparentnosti podľa článku 50 (neodložené); formálne právomoci Komisie pokutovať GPAI (článok 101)
2. decembra 2026	Nový zákaz AI generovaných intímnych zobrazení bez súhlasu a materiálu zobrazujúceho zneužívanie detí podľa článku 5 (pridaný omnibusom); koniec odkladu označovania pre generatívne systémy, ktoré už boli na trhu
2. augusta 2027	Povinnosť národných experimentálnych prostredí (posunutá z roku 2026 omnibusom); experimentálne prostredie na úrovni EÚ pri úrade pre AI
2. decembra 2027	Povinnosti pre vysokorizikové systémy z prílohy III (posunuté z augusta 2026 omnibusom)
2. augusta 2028	Povinnosti pre vysokorizikové systémy z prílohy I, teda bezpečnostné komponenty regulovaných výrobkov (posunuté z roku 2027 omnibusom)

Čítajte tú tabuľku zo svojho vlastného kresla. Ak nasadzujete chatbota, váš dátum je august 2026. Ak ste banka s úverovým modelom, december 2027 je čas, keď sa plná mašinéria vysokého rizika stane voči vám vymáhateľnou, čo znie ďaleko a nie je, vzhľadom na to, čo od vás kapitoly o požiadavkách žiadali. A ak vaša AI žije vnútri stroja alebo zdravotníckej pomôcky, máte čas do augusta 2028, s odvetvovými pravidlami, ktoré už dodržiavate, ako oporou medzitým.

Overte si

1. Niektó podá sťažnosť, že vaša firma nasadila systém na rozpoznávanie emócií na zamestnancov. Ktorý orgán vyšetroje a do ktorej úrovne pokút porušenie spadá? A) Úrad pre AI; do 15 miliónov eur alebo 3 percent obratu. B) Vnútroštátny orgán dohľadu nad trhom; do 35 miliónov eur alebo 7 percent obratu. C) Národné experimentálne prostredie; do 7,5 milióna eur alebo 1 percenta obratu. D) Európska rada pre AI; úroveň závisí od členského štátu.

2. Startup staviajúci vysokorizikový systém z prílohy III sa pýta, či musí pred vstupom na trh EÚ prejsť regulačným experimentálnym prostredím pre AI. Správna odpoveď je: A) Áno. Schválenie v experimentálnom prostredí je predpokladom prístupu na trh pre všetku vysokorizikovú AI. B) Áno, ale len kým sa neotvorí experimentálne prostredie na úrovni EÚ. C) Nie. Experimentálne prostredia sú dobrovoľná podpora; prístup na trh vedie cez posudzovanie zhody podľa článku 43 a výstupná správa z experimentálneho prostredia sa iba pozitívne zohľadní. D) Nie. Experimentálne prostredia sú otvorené len nasadzujúcim subjektom.

3. Kto spravuje databázu EÚ pre vysokorizikové systémy AI a kto do nej zapisuje systém poskytovateľa? A) Spravujú ju národné experimentálne prostredia a zapisujú systémy po schválení. B) Spravuje ju úrad pre AI; systémy zapisujú vnútroštátne orgány. C) Spravuje ju Komisia (článok 71); poskytovateľ svoj systém registruje priamo (článok 49). D) Každý členský štát spravuje vlastnú databázu; systémy zapisujú notifikované osoby.

4. Ktoré tvrdenie o vymáhaní AI Actu je k augustu 2026 presné? A) Komisia už podľa článku 101 pokutovala viacero poskytovateľov GPAI. B) Talianska pokuta pre chatbota Replika bola prvou pokutou podľa AI Actu. C) Nikde neexistuje potvrdená pokuta podľa AI Actu a len približne 8 z 27 členských štátov určilo svoje príslušné orgány načas. D) Vymáhanie sa nemôže začať, kým nebudú uverejnené harmonizované normy.

5. Po novele omnibus z roku 2026, odkedy sa uplatňujú povinnosti pre vysokorizikové systémy z prílohy III a aká je povaha toho dátumu? A) 2. augusta 2026, nezmenené oproti pôvodnému textu. B) 2. decembra 2027, pevný kalendárny dátum, nepodmienený dostupnosťou noriem. C) 2. decembra 2027, ale až

keď budú harmonizované normy citované v Úradnom vestníku. D) 2. augusta 2028, rovnaký dátum ako pre systémy z prílohy I.

Odpovede: 1: B. Odvodzovanie emócií na pracovisku porušuje článok 5 ods. 1 písm. f), zakázanú praktiku vymáhanú vnútroštátnym orgánom dohľadu nad trhom v najvyššej úrovni podľa článku 99 ods. 3. 2: C. Článok 57 ods. 5 robí z experimentálnych prostredí dobrovoľnú podporu inovácií a článok 57 ods. 7 hovorí, že výstupné správy sa „pozitívne zohľadnia“ na urýchlenie, nie nahradenie posudzovania zhody. 3: C. Článok 71 dáva databázu do rúk Komisie a článok 49 robí z registrácie vlastnú povinnosť poskytovateľa; experimentálne prostredia nehrajú žiadnu rolu. 4: C. Úprimná výsledková tabuľa ukazuje nula potvrdených pokút podľa AI Actu a sklz národného určovania orgánov, a pokuty Replika a Clearview boli prípady GDPR. 5: B. Omnibus posunul prílohu III na 2. december 2027 ako pevný dátum po tom, ako zákonodarcovia zamietli podmienenie termínov normami.

Teraz poznáte pravidlá, rozhodcov aj kalendár. Záverečná kapitola to všetko dá dokopy a pýta sa, čo rozumná organizácia naozaj urobí v pondelok ráno.

AI Act v praxi

Pamätáte na pondelkové ráno z kapitoly 3? Váš generálny riaditeľ vám prepošle článok o AI Acte s jednoriadkovou poznámkou: „Týka sa nás to?“ O deväť kapitol neskôr viete, že úprimná odpoveď znie „takmer určite, niekde“. Takže poznámka sa mení: „Asi by sme s tým mali niečo urobiť. Zoberiete si to?“ Teraz vlastníte súlad s AI Actom vo svojej firme. Kde vlastne začať?

Táto kapitola je moja odpoveď na tú otázku: rada, ktorú dávam, keď sa ma firma spýta „pochopili sme, že musíme byť v súlade, čo urobíme ako prvé?“. Potom na rozlúčku ako bonus vyriešime otázku, ktorú dostávam neustále: ak napíšete knihu alebo navrhnete kampaň s AI, núti vás Akt označiť to a vlastniť to vôbec? Odpoveď na obe časti je zaujímavejšia, než by ste čakali.

Všetko tu stavia na predchádzajúcich kapitolách. Nič tu nie je nové právo; je to ten istý Akt, videný zo strany človeka, ktorý s ním musí do štvrtka niečo urobiť.

Krok jeden: skontrolujte sa voči zákazom

Nie pravidlá pre vysoké riziko. Nie šablóny dokumentácie. Zákazy.

Tu je prečo. Zákazy podľa článku 5 sa uplatňujú od 2. februára 2025 (neprichádzajú, prišli) a sedia v najvyššej úrovni pokút, až 35 miliónov eur alebo 7 percent celosvetového obratu, so samotným režimom sankcií živým od augusta 2025. K augustu 2026 stále platí, že podľa Aktu zatiaľ nikto nikde pokutu nedostal. Ale ak je jedno miesto, kde by som nechcel byť prvý, je to tu.

Je tu aj skutočne blízky termín, o ktorom treba vedieť: nový zákaz AI generovaných intímnych zobrazení bez súhlasu a materiálu zobrazujúceho zneužívanie detí, „vyliekací“ zákaz pridaný do článku 5 novelou omnibus z roku 2026, sa začne uplatňovať 2. decembra 2026. A veľká vlna vysokého rizika pre prípady použitia z prílohy III pristane 2. decembra 2027. To je váš skutočný kalendár naliehavosti: december 2026 pre nový zákaz, december 2027 pre vysoké riziko. Nie panika, ale ani nekonečná dráha.

Cvičenie je teda jednoduché na opis a neprijemné na vykonanie: prejdite svoje produktívne prípady použitia AI a spýtajte sa, či sa niektorý z nich obtiera o osem (čoskoro deväť) zákazov. Pamätajte na recepčnú pod náladovou kamerou: odvodzovanie emócií na pracovisku je zakázané a je to presne ten druh funkcie „analytiky pohody zamestnancov“, ktorá potichu prichádza vnútri HR platformy. Ak nájdete podozrivý prípad použitia, berte ho vážne skoro. Vypnúť produktívny systém, ktorý zarába peniaze, je organizačne bolestivé a trvá to oveľa dlhšie, než ktokoľvek čaká. Lepšie to zistiť v pokojnom kvartáli než v liste od regulátora.

Keď už ste pri tom, urobte informačné stretnutie pre svojich dátových a AI ľudí. Povinnosť gramotnosti v oblasti AI podľa Aktu (článok 4) sa uplatňuje od toho istého februárového dátumu 2025. Omnibus jej znenie zmäkčil, ale prijímanie opatrení na budovanie gramotnosti sa stále očakáva a úprimne je to najlacnejšie opatrenie súladu, aké kedy kúpite.

Krok dva: zmapujte svoje roly a dávajte pozor na padacie dvere

Keď sú zákazy vyriešené, obídte organizáciu a inventarizujte všetko, čo by mohlo spadať pod definíciu systému AI. Pri každom odpovedzte na otázku, ktorá riadi väčšinu vašich povinností: sme tu poskytovateľ, alebo nasadzujúci subjekt?

Intuitívny test, ktorý poznáte z predchádzajúcich kapitol, stále robí väčšinu práce. Vyvinuli sme tento systém? Potom sme pravdepodobne poskytovateľ a Akt nás reguluje silno. Kúpili alebo prenajali sme si ho od niekoho iného? Potom sme pravdepodobne nasadzujúci subjekt s ľahšou sadou povinností.

Ale chcem, aby ste si do inventára niesli dve ďalšie cesty, lebo presne tam sa firmy klasifikujú nesprávne.

Po prvé, nemusíte nič trénovať, aby ste boli poskytovateľom. Článok 3 bod 3 hovorí, že poskytovateľom ste aj vtedy, keď si systém *dáte vyvinúť* niekým iným a uvediete ho na trh pod vlastným menom alebo ochrannou známkou. Ak vám agentúra postaví chatbota pre zákazníkov a ten sa dodáva ako „Asistent VašejFirmy“, fakt, že ste nikdy nenapísali riadok kódu, z vás nasadzujúci subjekt nerobí. Rozhoduje odznak na produkte.

Po druhé, nasadzujúci subjekt sa môže poskytovateľom *stať* dodatočne. Pri vysokorizikových systémoch článok 25 ods. 1 vymenúva troje padacích dverí:

AKT HOVORÍ · ČLÁNOK 25 ODS. 1

Každý distribútor, dovozca, nasadzujúci subjekt alebo iná tretia strana sa na účely tohto nariadenia považuje za poskytovateľa vysokorizikového systému AI... ak umiestnia svoje meno alebo ochrannú známku na vysokorizikový systém AI, ktorý už bol uvedený na trh; vykonajú jeho podstatnú zmenu; alebo menia zamýšľaný účel systému AI... tak, že sa stáva vysokorizikovým.

Prekročte ktorékoľvek z nich a zdedíte plné povinnosti poskytovateľa podľa článku 16; podľa článku 25 ods. 2 pôvodný poskytovateľ pri tom systéme z obrazu do veľkej miery vystupuje. Takže firma, ktorá kúpi HR nástroj na triedenie, preznačuje ho a model silno prispôsobí, nie je „len nasadzujúci subjekt“, nech faktúra hovorí čokoľvek.

Praktické ponaučenie: pri každom systéme v inventári zaznamenajte nielen „poskytovateľ alebo nasadzujúci subjekt“, ale aj „pod čím menom beží a upravili sme ho alebo sme mu zmenili účel?“. Verte mi: vaše budúce ja poďakuje vášmu minulému ja za toto mapovanie.

Krok tri: dokumentujte priebežne

Akt spomína dokumentáciu takmer na každom kroku: technickú dokumentáciu, posúdenia, hodnotenia, logy. Ale moja rada ide nad rámec formálnych požiadaviek: veďte si aspoň neformálny záznam o samotnej ceste k súladu. Ktoré systémy ste našli, čo ste o každom rozhodli, prečo ste niečo klasifikovali ako mimo rozsahu, kedy ste urobili to informačné stretnutie.

Ak budete niekedy potrebovať regulátorovi (alebo vlastnému predstavenstvu) ukázať, že ste konali dôsledne, „urobili sme to“ bez papierovej stopy má veľmi malú hodnotu. Datovaný priečinok nedokonalých poznámok porazí dokonalú pamäť zakaždým.

POZOR

Certifikát ISO/IEC 42001 nie je súlad s AI Actom. Užitočný ako chrbtica riadenia, ale k augustu 2026 nebola v Úradnom vestníku citovaná žiadna harmonizovaná norma, takže žiadny certifikát nekupuje predpoklad zhody.

Krok štyri: sledujte svojho regulátora a používajte Service Desk

Tu musím byť úprimný o stave sveta: pre mnohých z vás „váš národný regulátor“ ešte plne neexistuje.

STAV K AUGUSTU 2026

Len približne osem z dvadsiatich siedmich členských štátov určilo svoje vnútroštátne orgány načas. Na Slovensku a v Česku, kde sedí mnoho mojich študentov, sú vykonávacie zákony stále návrhmi: slovenský (MIRRI ako orgán dohľadu nad trhom) prešiel 14. augusta 2026 prvým vládny kolom, český (ČTÚ) je stále v legislatívnom procese.

To nie je dôvod uvoľniť sa; je to dôvod sledovať, lebo keď váš orgán naozaj vznikne, chcete byť medzi prvými, kto pozná jeho meno.

Dve konkrétne veci, ktoré môžete urobiť dnes. Po prvé, uložte si do záložiek AI Act Service Desk Komisie (ai-act-service-desk.ec.europa.eu), živý od októbra 2025. Obsahuje kontrolór súladu (Compliance Checker) a prehliadač Aktu (AI Act Explorer) a je to najbližšie k oficiálnej linke pomoci, akú Akt momentálne má. Po druhé, držte si na radare regulačné experimentálne prostredia. Povinnosť každého členského štátu prevádzkovať aspoň jedno posunula novela omnibus z roku 2026 na august 2027 a popri nej pridala experimentálne prostredie na úrovni EÚ pri úrade pre AI.

A tu je tip, ktorý môžete pokojne ignorovať: ak je vaša organizácia silno exponovaná, s mnohými prípadmi použitia pravdepodobne vysokorizikovými, zvážte pripraviť si už teraz jeden či dva *cvičné* prípady použitia. Nie systémy, ktoré nutne plánujete nasadiť; rozpracované, vierohodné návrhy, možno jeden s nízkym rizikom a jeden s vysokým. Keď sa blízko vás otvorí experimentálne prostredie, máte hotový, legitímny dôvod

prihlásiť sa a naučiť sa proces zvnútra. Vyhrávajú obe strany: vy získate kontext a vzťah, experimentálne prostredie skorého pokusného králik. Radšej by som bol firmou, ktorú regulátor už pozná, než tou, ktorú stretne prvýkrát pri inšpekcii.

Krok päť: viacero ľudí, nie jeden

Vzorec, ktorý som sledoval viac než raz: firma vymenuje presne jedného človeka, aby AI súlad „vlastnil“. Ten človek strávi mesiace tým, že sa stane skutočne znalým. A potom ho nájde headhunter, lebo vyškolených ľudí na AI súlad je málo a ubúda ich, a celé vaše inštitucionálne poznanie odíde dverami s 30-percentným zvýšením platu.

Takže: vyhradte viacero ľudí, z oboch strán organizácie, staviteľskej aj audítorskej. Redundancia tu nie je byrokratická výplň; je to spôsob, ako nechať poznanie prežiť kontakt s trhom práce.

MÔJ POHĽAD

Nebudte defenzívni. „Vyhne sa AI Actu tým, že nebudeme používať AI“ je stratégia, ktorú som naozaj počul, a je to spôsob, ako náklady na súlad zaplatiť aj tak: v stratenej konkurenčnej výhode namiesto právnických hodín. Akt je sada mantinelov, nie múr.

Bonus: AI, autorské právo a kto vlastní vašu knihu

A teraz otázka za polovicou e-mailov, ktoré dostávam. Ste autor používajúci textový model na pomoc pri písaní knihy alebo marketér generujúci plagát pre kampaň. Dve obavy, zvyčajne zamotané dokopy: musím to označiť ako vyrobené AI a vlastním to vôbec?

Rozmotajme ich, lebo Akt odpovedá na prvú a od druhej sa zámerne drží bokom.

Musíte označovať prácu s pomocou AI?

Tu je časť, ktorá ľudí prekvapí: v Akte neexistuje žiadna všeobecná povinnosť označovať prácu, ktorú ste vytvorili s pomocou AI. Povinnosti v oblasti transparentnosti podľa

článku 50, ktoré sa uplatňujú od 2. augusta 2026 a omnibus ich neodložil, sú konkrétne a väčšinou na vás ako autora či marketéra nedopadnú.

Povinnosť strojovo čitateľne označovať syntetický obsah sedí na *poskytovateľovi* generatívneho systému, nie na vás ako jeho používateľovi (článok 50 ods. 2). A aj tá povinnosť má výnimku, ktorú sa oplatí citovať, lebo kreslí presne tú čiaru, o ktorú sa ľudia obávajú.

AKT HOVORÍ · ČLÁNOK 50 ODS. 2

Táto povinnosť sa neuplatňuje v rozsahu, v ktorom systémy AI vykonávajú pomocnú funkciu pri štandardných redakčných úpravách alebo pokiaľ podstatne nemenia vstupné údaje poskytnuté nasadzujúcim subjektom alebo ich sémantiku...

Zrozumiteľne: korektorský koniec spektra je výslovne vonku. Ak AI vyleští vašu gramatiku a podstatne nezmení, čo ste napísali ani čo to znamená, žiadna povinnosť označovania nevzniká. To odpovedá aj na obavu „robí každý retuš vo Photoshope z môjho obrázka AI vygenerovaný?“.

Ako nasadzujúci subjekt dlhujete zverejnenie presne v dvoch situáciách (článok 50 ods. 4). Jedna: deep fake, obsah zobrazujúci skutočné osoby, miesta alebo udalosti spôsobom, ktorý by sa falošne javil ako autentický, sa musí zverejniť, so zmäkčenou, nerušivou formou zverejnenia pri zjavne umeleckých, kreatívnych alebo satirických dielach. Dva: AI vygenerovaný text uverejnený s cieľom informovať verejnosť o záležitostiach verejného záujmu (myslite správy) sa musí zverejniť, *pokiaľ* text neprešiel ľudskou alebo redakčnou kontrolou a niekto zaň nenesie redakčnú zodpovednosť. Tá posledná výnimka je v podstate redakčná výnimka: redaktor, ktorý za textom stojí, nahrádza štítok.

Prežňte cez to svoje dva scenáre. Román s pomocou AI nie je ani deep fake, ani informácia verejného záujmu, takže žiadna povinnosť nasadzujúceho subjektu. Bežný marketingový vizuál, žiadni falšovaní skutoční ľudia, tiež žiadna povinnosť nasadzujúceho subjektu. Štítky, ktoré čoraz častejšie vidíte na platformách, pochádzajú z povinnosti označovania na strane poskytovateľov a z politik platforiem, nie z povinnosti na vás.

Čo Akt naozaj hovorí o autorskom práve

Ľudia niekedy očakávajú, že AI Act je zákon o autorskom práve. Nie je. Ale autorské právo reguluje v jednom konkrétnom bode pipeline: na *ustupnej* strane, kde sa modely trénujú na dielach iných ľudí.

Každý poskytovateľ modelu AI na všeobecné účely musí zaviesť politiku súladu s autorským právom EÚ, a najmä identifikovať a rešpektovať strojovo čitateľné „výhrady“ (opt-out), ktoré nositelia práv môžu vyhlásiť podľa článku 4 ods. 3 smernice o autorskom práve z roku 2019, výhrady práv voči vyťažovaniu textov a dát (článok 53 ods. 1 písm. c)). Navrch musí každý poskytovateľ GPAI zverejniť dostatočne podrobné zhrnutie obsahu použitého na tréňovanie podľa šablóny úradu pre AI, ktorá je povinnou šablónou od decembra 2025 (článok 53 ods. 1 písm. d)). Kódex postupov pre GPAI má dokonca samostatnú kapitolu o autorskom práve, ktorá poskytovateľov tým prevedie. To všetko je účinné od augusta 2025 a omnibus sa toho nedotkol.

Takže ak bol váš román zoškrabaný do tréningovej množiny napriek vašej vyhlásenej výhrade, AI Act je veľmi váš zákon. Ak sa pýtate, kto vlastní plagát, ktorý pre vás model vygeneroval, nie je.

Takže kto vlastní výstup?

Pre vlastníctvo musíme AI Act úplne opustiť a siahnuť po autorskom práve, čo v EÚ ako palcové pravidlo znamená národný autorský zákon krajiny, kde tvoríte. Prešiel som si tým osobne: patril som medzi prvých ľudí na Slovensku, ktorí poriadne publikovali knihu napísanú generatívnymi modelmi, text aj obrázky. Tak som si sadol so svojím právnikom a nakreslili sme tri scenáre.

Scenár jeden: AI je *štetec*. Štetec je statický nástroj. Urobte identický pohyb rukou dvakrát a namaľujete identický ťah dvakrát. Zo samotného nástroja neplynie žiadna jedinečnosť.

Scenár dva: AI je *pes držiaci štetec*. Povedzte psovi, aby dvakrát namaľoval kruh, a dostanete dva rôzne kruhy; nedokonalosť jeho svalovej kontroly niečo pridá. Teraz je tu jedinečnosť, ale pes je stále váš nástroj; vy ste ho vycvičili a povedali mu, čo maľovať.

Scenár tri: AI je *tvorivá entita* sama osebe, v ktorom prípade by práva k tomu, čo spolu vytvoríte, držala ona, nie vy.

Ktorý to je? Slovenský autorský zákon (a autorské zákony sú naprieč EÚ silno harmonizované, takže ten váš bude znieť podobne) definuje predmet autorského práva ako dielo z oblasti literatúry, umenia alebo vedy, „ktoré je jedinečným výsledkom tvorivej duševnej činnosti autora vnímateľným zmyslami“. Nosné slová sú *jedinečný výsledok a tvorivá duševná činnosť*.

Najprv jedinečnosť. Položte textovému modelu tú istú otázku v dvoch čerstvých konverzáciách a dostanete dve rôzne odpovede, lebo vývojári do generovania zámerne zabudovávajú náhodnosť: model vzorkuje medzi najpravdepodobnejšími ďalšími slovami namiesto toho, aby vždy vybral to najvyššie. Takže scenár so štetcom odpadá: tento nástroj sa neopakuje.

Ostáva pes a tvorivá entita. A tu rozhoduje princíp starší než akákoľvek regulácia AI: v autorskom práve, ako stojí, medzinárodne, sa za schopných tvorivej činnosti považujú len ľudia. Stroje nie sú autormi. Čo znamená, že dnešná generatívna AI pristáva priamo v scenári dva: veľmi talentovaný pes. Produkuje jedinečné výstupy, ale tvorivá duševná činnosť, a teda autorské právo, patrí človeku, ktorý ju riadil. Právne sedí váš textový model na tej istej poličke ako Microsoft Word: nástroj, nezvyčajne živý.

Obvyklé upozornenie tu platí s extra silou: toto je súčasné usporiadanie, národné zákony sa v detailoch líšia a súdy okraje ešte testujú. Koľko ľudského riadenia je dosť, je presne ten druh otázky, ktorý v najbližších rokoch prinesie zaujímavé rozsudky. Ak na vašej odpovedi visia skutočné peniaze, toto je rozhovor, ktorý treba viesť s vlastným právnikom, vo vlastnej krajine. Ale ako pracovný model mi pes so štetcom (aj mojej knihe) poslúžil dobre.

Overte si

1. Váš generálny riaditeľ vám dnes odovzdá súlad s AI Actom. Podľa tejto kapitoly je váš prvý ťah: A) Z začať písať technickú dokumentáciu pre každý systém AI B) Skontrolovať svoje produktívne prípady použitia AI voči zákazom podľa článku 5, ktoré sa uplatňujú od februára 2025 C) Počkať, kým bude určený váš národný regulátor D) Požiadajte o certifikáciu ISO/IEC 42001, ktorá sa počíta ako súlad s AI Actom

2. Firma kúpi vysokorizikový systém na triedenie uchádzačov, preznačuje ho pod vlastnú ochrannú známku a model podstatne upraví. Podľa článku

25 ods. 1 firma: A) Ostáva nasadzujúcim subjektom, keďže systém nevyvinula B) Stáva sa poskytovateľom a preberá povinnosti poskytovateľa podľa článku 16 C) Delí sa o povinnosti poskytovateľa 50/50 s pôvodným dodávateľom D) Musí len informovať pôvodného poskytovateľa

3. Kde AI Act reguluje autorské právo? A) Prideluje autorské práva k AI výstupom poskytovateľovi modelu B) Obsahuje úplný režim autorského práva EÚ pre diela vygenerované AI C) Na strane tréningových vstupov: poskytovatelia GPAI potrebujú politiku autorského práva rešpektujúcu výhrady voči vyťažovaniu textov a dát a musia zverejniť zhrnutie tréningového obsahu D) Nikde, lebo Akt autorské právo nikdy nespomína

4. Spisovateľka použije AI model ako korektora, ktorý urobí minimálne úpravy, a potom knihu vydá. Podľa článku 50: A) Kniha musí niesť štítok „vygenerované AI“ B) Spisovateľka musí knihu zaregistrovať v databáze EÚ C) Nevzniká žiadna povinnosť, keďže pomocné úpravy, ktoré podstatne nemenia obsah, sú vyňaté, a román nie je ani deep fake, ani informácia verejného záujmu D) Vydavateľ musí každú stranu strojovo čitateľne opatriť vodoznakom

5. Kto podľa súčasného autorského práva drží autorské práva k obrázku, ktorý vygenerujete pomocou AI nástroja? A) AI model, keďže vyprodukoval jedinečný výsledok B) Nikto, lebo diela s pomocou AI sa nedajú vôbec chrániť autorským právom C) Vývojár modelu, cez podmienky používania D) Vy, človek, ktorý nástroj riadil, keďže za schopných tvorivej činnosti sa považujú len ľudia

Odpovede. 1: B. Zákazy sú najstaršie živé povinnosti, nesú najvyššiu úroveň pokút a odvinutie zakázaného prípadu použitia trvá, tak idú prvé. 2: B. Preznačkovanie, podstatná zmena alebo zmena účelu na vysokorizikový robia z nadväzujúceho aktéra poskytovateľa podľa článku 25 ods. 1, pričom pôvodný poskytovateľ podľa ods. 2 ustupuje. 3: C. Článok 53 ods. 1 písm. c) a d) reguluje stranu tréningových vstupov, kým vlastníctvo výstupov je ponechané národnému autorskému právu. 4: C. Povinnosť označovania podľa článku 50 ods. 2 sedí na poskytovateľoch a vylučuje pomocné štandardné úpravy a povinnosti zverejnenia nasadzujúceho subjektu v článku 50 ods. 4

pokrývajú len deep fake a text verejného záujmu. 5: D. Dnešné právo berie generatívnu AI ako nástroj (pes so štetcom): jej výstupy sú jedinečné, ale tvorivá činnosť, a teda autorské právo, je vyhradená ľuďom.

A to je celá cesta. Začali sme definíciou, ktorá zdanlivo hltala každú tabuľku v Európe, a končíme tým, že vlastníte knihu, ktorú vám pomohol napísať váš veľmi talentovaný pes. Cestou ste sa naučili rozpoznať zákaz, klasifikovať vysokorizikový systém, rozlíšiť poskytovateľa od nasadzujúceho subjektu a prečítať termín súladu bez paniky. To je úprimne viac, než dokáže väčšina ľudí radiacich k tejto regulácii.

Akt sa bude ďalej hýbať: dátumy sa už raz posunuli, usmernenia stále prichádzajú a jedného dňa bude aj prvá pokuta, o ktorej budú noviny písať. Keď sa to stane, neberte nikomu za slovo, čo to znamená, vrátane mňa. Otvorte text. Teraz sa v ňom vyznáte.

Ďakujem za prečítanie. Teraz choďte zmapovať svoje systémy a možno nechajte psa, nech vám cestou niečo namaľuje.



TRANSPARENTNOSŤ AI

Táto publikácia vznikla s pomocou generatívnej AI (Claude od spoločnosti Anthropic) v spolupráci s Robertom Barcikom. Text prešiel jeho revíziou a on nesie redakčnú zodpovednosť za to, čo je tu zverejnené (LearningDoe s.r.o.). Tam, kde AI hovorí vlastným hlasom, text to uvádza. Zverejnené v duchu článku 50 Aktu EÚ o umelej inteligencii.