

State of AI Dev

Hol tart a magyar szoftverfejlesztés AI-native átalakulása? Mennyit ír már a gép, mit vett át az agent, ki ellenőrzi, és mennyi fedezet van mögötte. Fejlesztők és mérnöki vezetők válaszai alapján.

65%

minden nap AI-eszközzel dolgozik

39%

naponta delegál autonóm agentnek

48%

szerint a kódja fele már AI-generált

Önkéntes, önbevallásos minta – nem reprezentatív a magyar fejlesztői populációra.

Tartalom

01 Vezetői összefoglaló

02 Kulcsállítások

03 Archetípusok

04 Agentic coding

05 AI-generált kód

06 Review & verifikáció

07 Governance

08 Biztonság

09 Tanulási igények

10 Módszertan

Forrás: online kérdőív. Minden arány a válaszadókra vonatkozik. Az adatok önbevallásosak, a minta nem reprezentatív.

01 – VEZETŐI ÖSSZEFOGLALÓ

Az AI-native fejlesztés megérkezett. A fedezete nem.

A minta nem a fejlesztői átlagot mutatja, hanem azt a réteget, amelyik már ma úgy dolgozik, ahogy a többség valószínűleg egy-két éven belül. Ebben a rétegben az AI-használat kérdése eldőlt — a nyitott kérdés az, hogy mi van alatta.

1 A magyar fejlesztői minta nagy része átlépett a kísérletezésen: **65%** minden nap AI-eszközzel dolgozik, és **48%** szerint a production kódjának több mint fele már AI-generált. Az AI nem eszköz a szerszámosládában, hanem a munkafolyamat alaprétege.

2 Az agentic coding a mintában már nem niche: **69%** nyúlt autonóm agenthez, **39%** naponta delegál end-to-end feladatot. Ez a réteg — az **AI-native csoport** — másképp dolgozik, mint a többi: más eszközöket használ, más munkafázisokat adott át, és másban látja az akadályt.

3 A verifikáció terhe nem együtt nő a használattal. A legnagyobb terhet a **középmezőny** viseli (3,11/5), az AI-native-ok a legkisebbet (2,49/5). Ez nem bizonyítja, hogy az AI könnyebbé teszi a review-t — inkább azt jelzi, hogy a rutin és a kialakult ellenőrzési szokások számítanak.

4 A legnagyobb szerkezeti probléma a **fedezetlen tempó**: az AI-érettség és a governance-érettség között gyakorlatilag nincs kapcsolat ($p=0,16$). A

startupok érettsége 68, governance-e 31; a nagyvállalatoknál fordított a helyzet (60 vs 67).

5

A biztonsági érettség a leggyengébb láncszem: **30%-**nál előfordul szenzitív adat publikus LLM-be juttatása, és az LLM-integrációt fejlesztők **70%-**ának nincs dedikált prompt-injection tesztje. A kockázat nem a lemaradóknál koncentrálódik, hanem a leggyorsabbaknál.

Az AI-native lánc: hol morzsolódik le a mezőny

egymásba ágyazott feltételek · minden lépés az előző szűkítése · a minta %-a



Minden sor tartalmazza az összes fölötte lévő feltételt is. A legalsó sor tehát azokat jelenti, akik napi szinten használnak AI-t, naponta delegálnak autonóm agentnek, és a kódjuk több mint fele AI-generált.

A tölcser első fele lapos, a második meredek: az **AI-használat** már alapállapot (93%), a **napi agent-delegálás** viszont még mindig szűrő (39%). A teljes AI-native láncot a minta **31%-a** teljesíti — ez a riport központi szegmense.

02 – KULCSÁLLÍTÁSOK

Tíz szám, ami leírja a 2026-os állapotot

Mindegyik állítás a lenti fejezetek valamelyikéből származik, és mindegyikhez tartozik ábra. A számok önbevalláson alapulnak.

37% Minden harmadik válaszadó AI-native

Ennyien adnak ki napi szinten end-to-end feladatot autonóm agentnek — nem asszisztensként használják az AI-t, hanem munkatársként. A teljes láncot (napi használat + napi agent-delegálás + 50%+ AI-kód) a minta **31%-a** teljesíti.

48% A production kód fele gépi kézből jön

Ennyien mondják, hogy a production kódjuk **több mint fele** AI-generált, emberi módosítás nélkül vagy minimálissal; **29%** a 71%+ sávban van. Ez önbevallás, nem repository-mérés — de a nagyságrend így is beszédes.

56% Az agentic coding senior gyakorlat

A **tech lead / architect** csoport delegál naponta a legnagyobb arányban, a senior fejlesztőknél 46% — a junioroknál viszont csak **12%**. Az agent-delegálás tapasztalatot és ítélőképességet igényel, nem váltja ki.

3,11 A verifikációs teher középen tetőzik

A review-teher nem a legmélyebb AI-használóknál a legnagyobb: a csúcs az asszisztens-használó és az agent-kísérletező csoport (mindkettő 3,11/5), a legalacsonyabb pedig épp az ai-native csoportnál (2,49/5). A középmezőny fizeti a tanulási görbe árát.

79% A minőség csak addig gát, amíg kevés az AI

Az AI-native-ok ennyien mondják, hogy **nincs akadály** a szélesebb használat előtt. Az asszisztens-használóknál viszont még **46%** a minőséget és megbízhatóságot jelöli fő blokkolónak.

0,16 Az AI-érettség és a governance nem jár együtt

A két index közti rangkorreláció gyakorlatilag nulla ($p=0,16$), az AI-kódarány és a governance-index között pedig **$p=0,02$** . Aki mélyen használja az AI-t, nem dolgozik szabályozottabb környezetben — csak gyorsabban.

19 pp Ekkora a rés az adoptáció és a fedezet között

Az adoptációs mutatók átlaga **57%**, a mögöttük álló governance-mutatóké **38%**. A cégek **75%-a** aktívan ösztönzi az AI-t, működő policy viszont csak **43%-nál** van.

30% A shadow AI nem a szélén van, hanem középen

Ennyien mondják, hogy előfordul szenzitív adat (forráskód, kulcs, séma) publikus LLM-be juttatása. Az AI-native csoportban a **rendszeres** szivárgás a leggyakoribb: **15%** — vagyis a legmélyebb használat jár a legnagyobb kitettséggel.

70% Aki LLM-et épít termékbe, jellemzően nem teszteli

Az LLM-integrációt fejlesztők közül ennyiennek nincs dedikált prompt-injection tesztje; mindössze **10%** red-teamel release előtt.

31% A tanulási igény együtt érik a használattal

Az AI-native-ok már **AI-alapú architektúrát** akarnak tanulni, az asszisztens-használók a **prompt engineeringet** (39%), a kívülállók pedig az **AI biztonságot és etikát** (54%). A képzési kínálatnak szintenként kell szólnia.

03 — ARCHETÍPUSOK & AI-ÉRETTSÉG

Négy fejlesztő, négy valóság

A „mennyire használ AI-t” kérdésre nincs egy válasz. A minta négy jól elkülönülő csoportra bomlik aszerint, hogy milyen **mélyen** engedik be az AI-t a munkafolyamatba — és ez a besorolás szinte minden más kérdésnél is megjósolja a választ.

● **15%**

Kívülálló

Ritkán vagy egyáltalán nem nyúl AI-hoz

legfeljebb havi néhány alkalommal használ AI-t

● **32%**

Asszisztens-használó

Heti/napi AI, de a kód az ő kezében marad

legalább hetente használ AI-t, de agentnek nem ad ki feladatot (vagy csak kísérletezik)

● **16%**

Agent-kísérletező

Heti szinten delegál kisebb taskokat

heti szinten ad ki kisebb taskokat autonóm agentnek

● **37%**

AI-native

Az agent a workflow szerves része

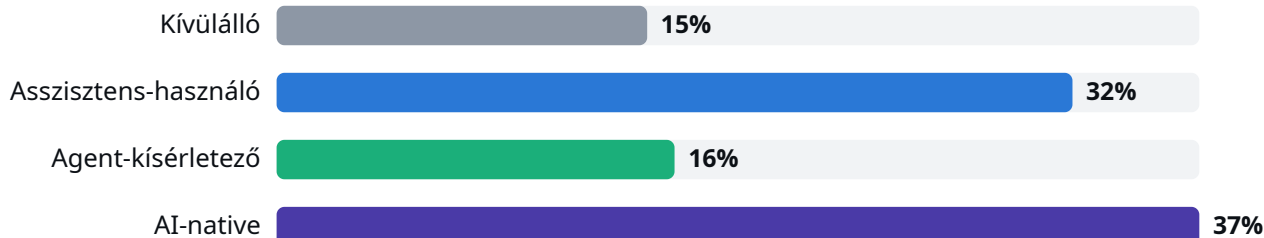
napi szinten delegál end-to-end feladatot autonóm agentnek

● ÍGY OLVASD – AZ ARCHETÍPUSOK

A besorolás **szabályalapú**, nem statisztikai klaszterezés: két kérdésből dől el (milyen gyakran használ AI-t, és milyen mélyen delegál autonóm agentnek), így bárki be tudja lőni, hova esne. A tengely a használat mélysége — nem képesség, nem senioritás és nem teljesítmény.

Az archetípusok megoszlása

a válaszadók %-a



A négy csoport profilja

Ugyanaz a tíz mutató mind a négy csoportra. Az árnyalat oszloponként normált: a sötétebb cella az adott mutató legmagasabb értéke a négy csoport közül.

Archetípus-profil

soronként egy csoport · oszloponként egy mutató · %-ok és 1–5 skálák

	Naponta használ AI-t	A kód 50%-a AI-generálta	Eszköz / fő	Review-teher (1-5)	Szerinte nincs akadály	Céges támogatás (1-5)	Van működő AI-policy	Kapott AI-képzést	Előfordul shadow AI	Nincs LLM-biztonsági teszt
Kívülálló	0%	12%	0,80	2,60	8%	3,10	27%	23%	27%	50%
Asszisztens-használó	59%	27%	2,00	3,10	23%	4,20	41%	43%	34%	59%
Agent-kísérletező	68%	29%	2,20	3,10	46%	4,50	54%	57%	25%	68%
AI-native	95%	89%	2,20	2,50	79%	4,80	46%	42%	31%	43%

A profil három dolgot mutat meg egyszerre. Egy: az **AI-kódarány** ugrásszerűen nő (27% → 89%), miközben az eszközök száma alig (2,00 → 2,20) — nem több eszközről van szó, hanem mélyebb bizalomról. Kettő: a **review-teher** nem a tetején a legnagyobb, hanem középen. Három: a **governance-mutatók** (policy, képzés, biztonsági teszt) nem követik az érettséget — az AI-native oszlopok itt nem világosabbak, mint a többi.

04 – AGENTIC CODING

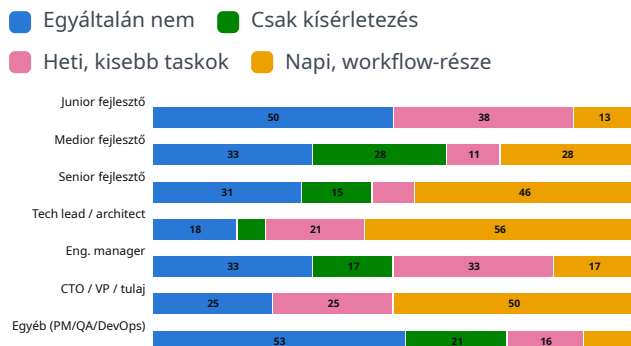
A delegálás senioritási kérdés

Az autonóm agent — ami maga elemzi a repót, megírja a változtatást és nyitja a PR-t — a mintában már nem demó. **69%** nyúlt hozzá, **39%** naponta használja. De hogy ki lép be, azt élesen szelektálja a szerepkör.

A napi agent-delegálás a **tech lead / architect** körében 56%, a senior fejlesztőknél 46% — a **junioroknál 12%**, a medioroknál 28%. Az agent kiadása nem a munka megspórolása: ítélőképesség kell hozzá, hogy valaki felmérje, mit lehet rábízni és mit nem.

Agent-autonómia szerepkörönként

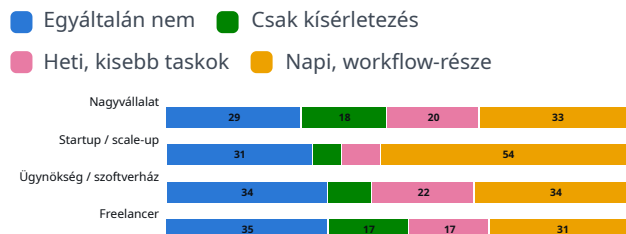
soronként 100% · q12b × q03



A kis szerepkör-cellák (junior, eng. manager) a trendet mutatják, a pontos %-ot ne.

Agent-autonómia cégtípusonként

soronként 100% · q12 × q03

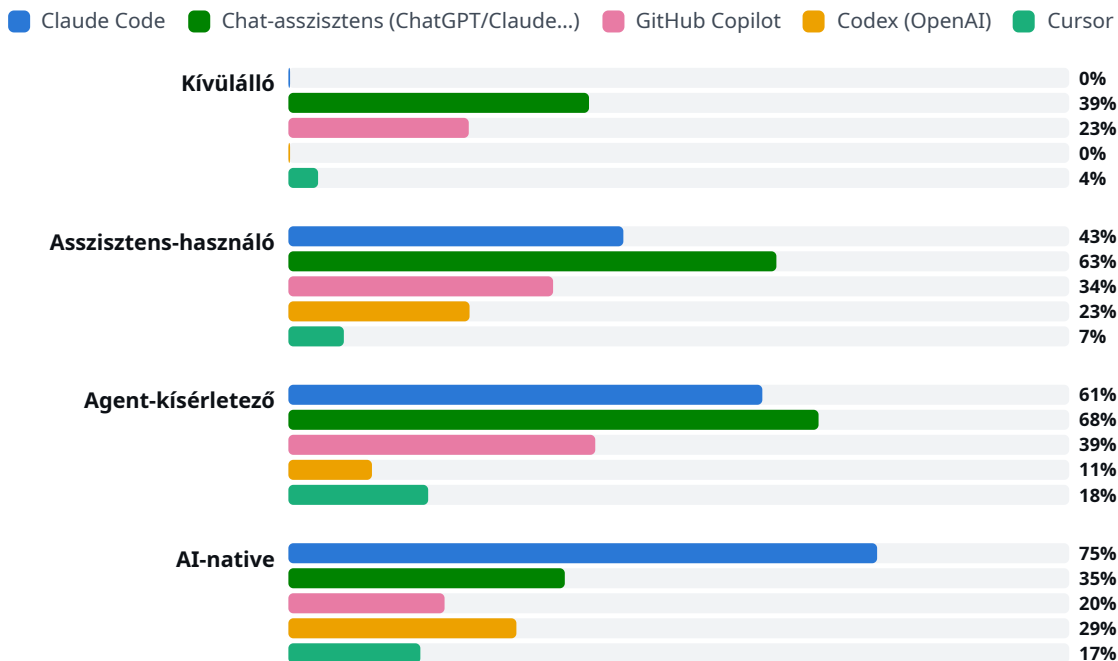


Az eszközkészlet is elválik

Nem ugyanazok az eszközök futnak a négy csoportnál. A chat-asszisztens a belépő szint jellemzője; ahogy nő a delegálás mélysége, átveszi a helyét a terminál- és repo-szintű agent.

Eszközhasználat archetípusonként

a csoport hány %-a használja az adott eszközt · több válasz lehetséges



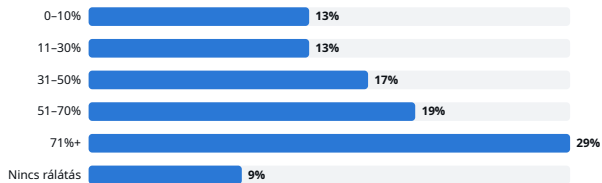
A **Claude Code** a minta egészében 51%-os, az AI-native csoportban viszont **75%**. A chat-asszisztens fordítva mozog: az asszisztens-használóknál 63%, az AI-native-oknál már csak 35%. Átlagosan 1,91 eszköz fut fejenként — az eszközök kiegészítik, nem kiváltják egymást.

Mennyit ír már tényleg a gép

A kérdés a production kód arányára vonatkozott, emberi módosítás nélkül vagy minimálissal. **48%** mondja, hogy ez több mint a fele, **29%** pedig a 71%+ sávba sorolja magát. Ez becslés, nem mérés — de a szegmensek közti különbség így is beszédes.

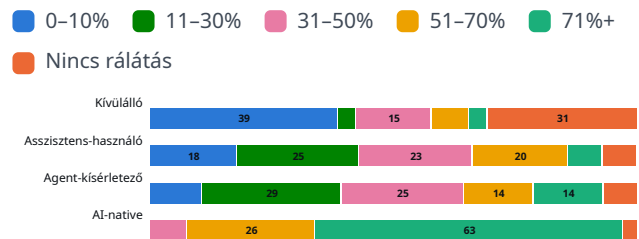
AI-generált kód aránya

q05 · a válaszadók megoszlása



AI-kódarány archetípusonként

soronként 100% · archetípus × q05

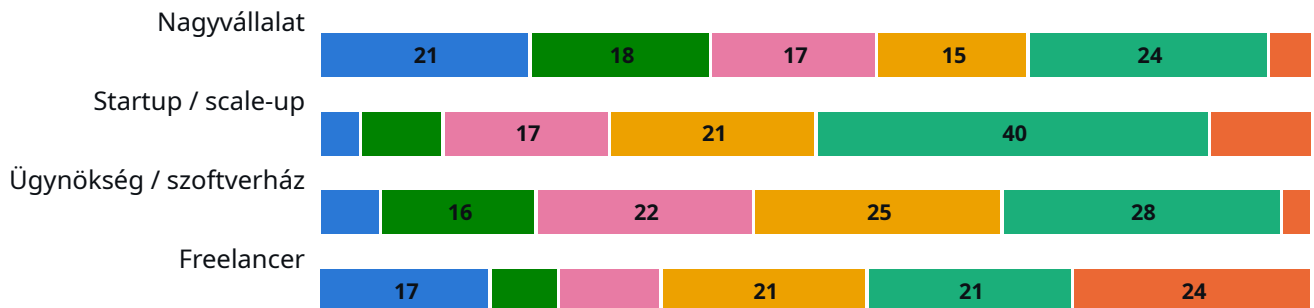


Az archetípus szinte determinálja a kódarányt: az AI-native csoport **89%**-ánál a kód több mint fele AI-generált, az asszisztens-használóknál 27%. Az agent-autonómia és az AI-kódarány között ez a minta legerősebb együttjárása (**p=0,67**) — aki komplett feladatokat delegál, annál a kód is nagyobb részben gépi.

AI-kódarány cégtípusonként

soronként 100% · q12 × q05

0–10% 11–30% 31–50% 51–70% 71%+ Nincs rálátás

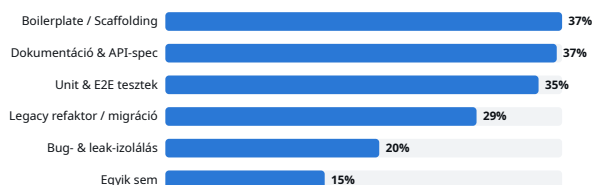


Melyik munkafázist vette át az AI

Először a jól körülhatárolt, olcsón ellenőrizhető munka ment át: boilerplate, tesztek, dokumentáció. A mély kontextust igénylő feladatok — legacy refaktor, bug-izolálás — csak a magasabb érettségi szinteken nyílnak meg.

Átvett munkafázisok

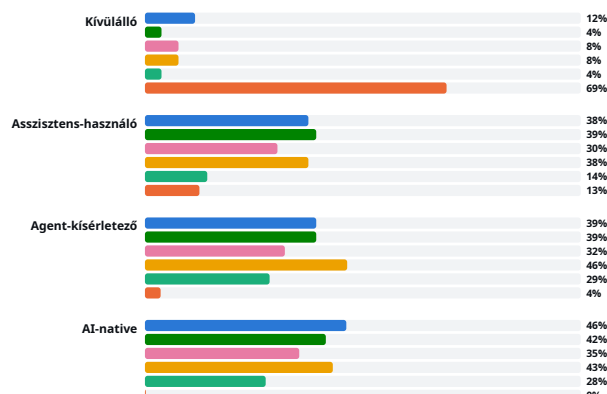
q04 · max. 2 válasz · a válaszadók %-a



Munkafázisok archetípusonként

a csoport hány %-a jelölte · több válasz lehetséges

- Boilerplate / Scaffolding
- Unit & E2E tesztek
- Legacy refaktor / migráció
- Dokumentáció & API-spec
- Bug- & leak-izolálás
- Egyik sem



A különbség nem a sorrendben van, hanem a **mélységben**. A bug- és leak-izolálást — a legnehezebben delegálható feladatot — az AI-native csoport **28%-a** jelölte, az asszisztens-használók 14%-a. A kívülállók 69%-a szerint az AI *egyik* fázisban sem váltott ki manuális munkát; az AI-native csoportban ezt **senki** nem jelölte.

06 – REVIEW & VERIFIKÁCIÓ

A verifikációs völgy

Kézenfekvő feltevés, hogy minél több az AI-kód, annál nagyobb az ellenőrzés terhe. Az adat nem ezt mutatja: a terhelés **a középmezőnyben tetőzik**, és épp a legmélyebb használóknál a legalacsonyabb.

Review-teher archetípusonként

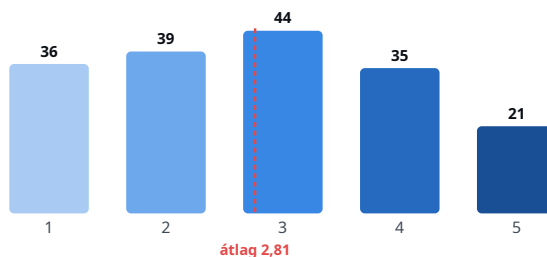
q06 átlaga · 1: nem több · 5: sok extra idő



Review-teher eloszlása

átlag 2,81/5

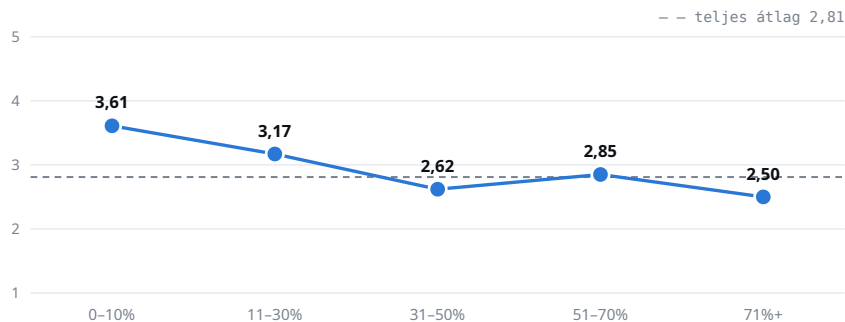
q06 · a teljes minta



A görbe alakja a lényeg: a csúcs **az asszisztens-használó és az agent-kísérletező csoport (mindkettő 3,11/5)**, a legalacsonyabb pedig **az ai-native** csoport (2,49/5) — 0,62 pont a különbség. Ugyanez látszik az AI-kódarány mentén is: a legmagasabb terhet a **0–10%-os** sáv jelzi (3,61/5), a legalacsonyabbat a 71%+ (2,50/5).

Review-teher az AI-kódarány mentén

q06 átlaga a q05 sávjai szerint · a „nincs rálátás” kihagyva



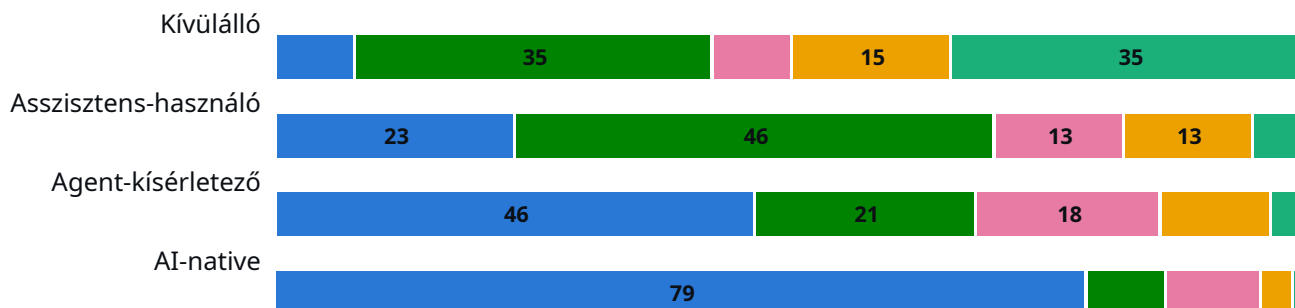
● AMIT EZ NEM JELENT

Ebből **nem** következik, hogy az AI-kód könnyebben review-zható. Legalább három magyarázat fér az adatra: (1) aki sokat használ, gyakorlottabban szűri a gépi hibákat; (2) aki nehéznek találta, abbahagyta vagy vissza sem lépett — vagyis a magas terhet érzők a kevés-AI sávokban maradtak; (3) aki nagyon mélyen delegál, esetleg *kevesebbet* ellenőriz, és ezt éli meg kisebb teherként. A kérdőív a **szubjektív terhet** mérte, nem a review minőségét vagy a hibák számát — a harmadik magyarázat kockázatát az adat nem zárja ki.

Mi a fő akadály a szélesebb használatban

soronként 100% · archetípus × q11

■ Nincs akadály ■ Minőség / megbízhatóság ■ Biztonság / adatvédelem ■ Hiányzó engedély / licenc
■ Szabályozás / EU AI Act



Az akadály **szintfüggő**. Az asszisztens-használóknál a minőség és megbízhatóság vezet (46%), az AI-native csoport **79%-a** szerint viszont már nincs akadály. A kívülállóknál a szabályozói megfelelés és a minőség egyforma súllyal áll — vagyis a belépést nem egy dolog gátolja.

07 – VÁLLALATI GOVERNANCE

Fedezetlen tempó

A cégek nem fékeznek: **75%** aktívan ösztönzi az AI-használatot (átlag 4,30/5). A szabályzat, a képzés és a technikai fedezet viszont nem tartott lépést — és a kettő között nincs érdemi kapcsolat.

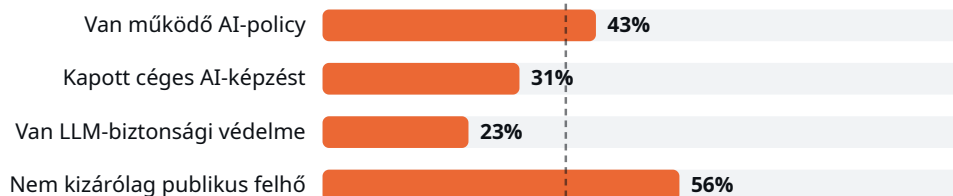
Az adoptáció és a fedezet közti rés

a minta %-a · felül a használat mélysége, alul a mögötte álló feltételek

ILYEN MÉLYEN HASZNÁLJUK



ENNYI VAN ALATTA



19 százalékpontos rés

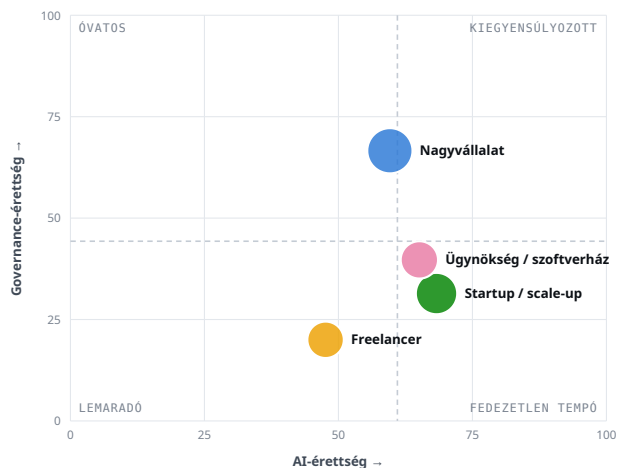
Az adoptációs mutatók átlaga **57%**, a governance-mutatóké **38%** — 19 százalékpont a különbség. A legélesebb egyedi kontraszt: 75% ösztönzés mellett 43% működő policy és 31% céges képzés.

Érettség vs. fedezet szegmensenként

Két 0–100-as index egymással szemben: mennyire mélyen használják az AI-t, és mennyire van mögötte szabályzat, képzés és biztonsági felkészültség. Ha a kettő együtt mozogna, a pontok átlóban állnának. Nem állnak.

AI-érettség × governance-érettség

cégtípusonként · a pont mérete a szegmens létszáma · a szaggatott vonal a minta átlaga

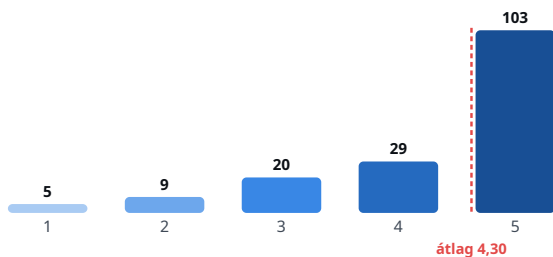


A **Startup / scale-up** szegmens érettsége 68, a governance-e 31 — a „fedezetlen tempó” kvadráns. A **Nagyvállalat** fordítva áll (60 vs 67): lassabban halad, de van mögötte keret. A két index közti rangkorreláció a teljes mintán **$p=0,16$** , az AI-kódarány és a governance-index között pedig **$p=0,02$** — statisztikailag ez azt jelenti, hogy **a kettőnek semmi köze egymáshoz**.

Céges támogatottság

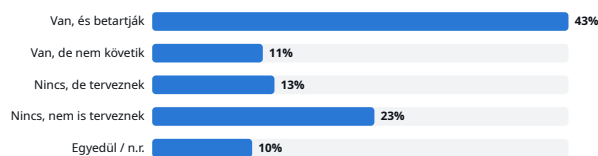
átlag 4,30/5

q07 · 1: tiltja · 5: aktívan ösztönzi



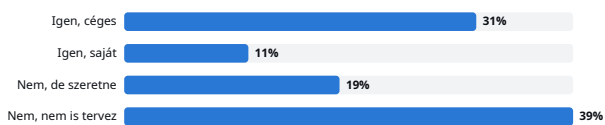
Van-e formális AI-policy?

q08 · szabályzat státusza



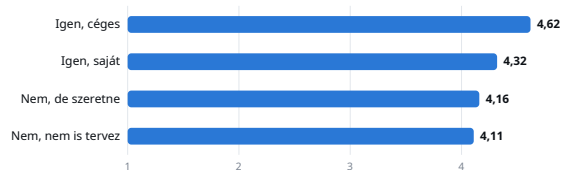
AI-képzés az elmúlt 12 hónapban

q13 · formális képzés



Céges támogatás a képzési státusz szerint

átlagos q07-érték (1–5) csoportonként



A képzés a leglátványosabb hiány: **58%** nem kapott formális AI-képzést az elmúlt egy évben, és csak 31% kapott céges keretből. Ez ugyanabban a mintában történik, ahol 65% minden nap AI-jal dolgozik — a készségfejlesztés a fejlesztők magánügye maradt.

08 — AI-BIZTONSÁG

A leggyorsabbak a legkitettebbek

Két külön kockázat fut egymás mellett: a **kifelé menő adat** (mi kerül be a promptba) és a **befelé jövő támadás** (mi történik, ha az LLM-integrációt valaki manipulálja). Egyikben sincs jó helyzetben a minta.

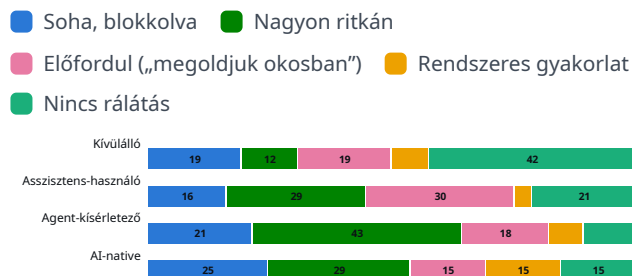
Szenzitív adat publikus LLM-be (shadow AI)

q09 · forráskód / kulcs / séma szivárgása



Shadow AI archetípusonként

soronként 100% · archetípus × q09

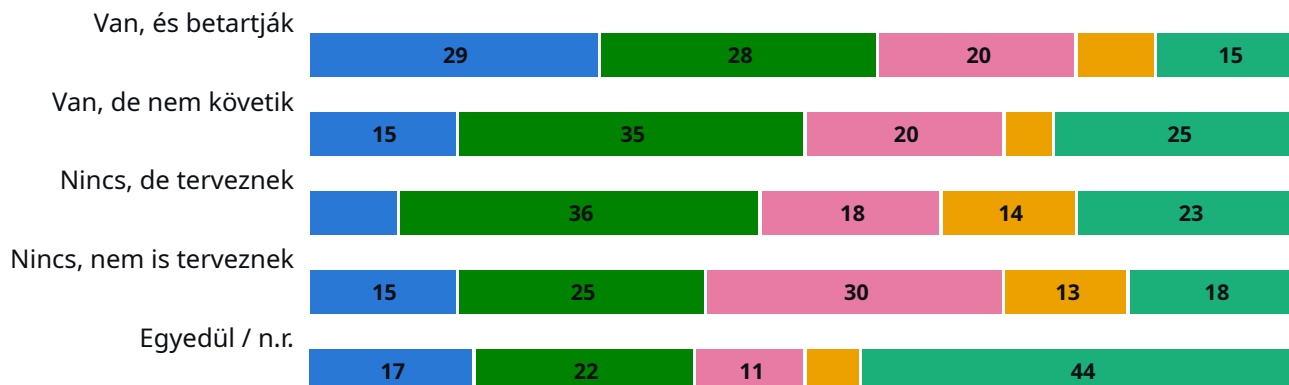


Összesen **30%**-nál fordul elő szenzitív adat publikus LLM-be juttatása. A megoszlás viszont ellentmond az intuíciónak: az AI-native csoportban a **rendszeres** gyakorlat a leggyakoribb (15%), és közülük 48% dolgozik működő policy nélkül — miközben ez a csoport küldi a legtöbb kódot modellbe. A kitettség nem a lemaradónál koncentrálódik.

Visszafogja-e a szabályzat a szivárgást?

soronként 100% · q08 × q09

■ Soha, blokkolva ■ Nagyon ritkán ■ Előfordul („megoldjuk okosban”) ■ Rendszeres gyakorlat ■ Nincs rálátás



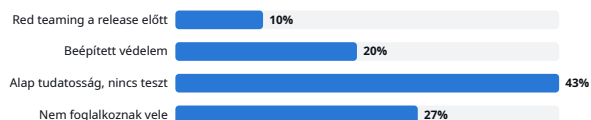
A policy segít, de nem old meg mindent. Ahol van szabályzat és betartják, ott **29%** mondja, hogy soha nem fordul elő — ahol nincs és nem is terveznek, ott csak 15%. De a működő policy mellett is 20% jelzi, hogy „megoldják okosban”.

LLM-integráció: aki épít, az véd?

A minta **76%**-a fejleszt LLM-integrációt. Náluk a prompt injection és a RAG-mérgezés nem elméleti kérdés, hanem termékkockázat.

Felkészültség LLM-sebezhetőségekre

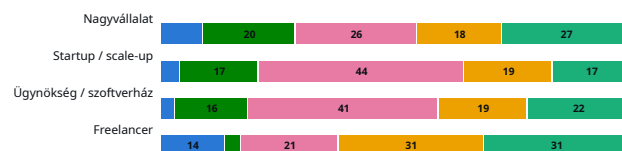
csak az llm-integrációt fejlesztők



Biztonsági érettség cégtípusonként

soronként 100% · q12 × q14 · a teljes minta

■ Red teaming a release előtt ■ Beépített védelem
■ Alap tudatosság, nincs teszt ■ Nem foglalkoznak vele
■ Nem fejleszt LLM-integrációt



Az LLM-integrációt fejlesztők **70%**-ának nincs dedikált tesztje, ezen belül 27% egyáltalán nem foglalkozik a kérdéssel. Release előtt **10%** red-teamel. Ez az a pont, ahol a riport számai a leginkább figyelemztető jelzések, nem büszkeségre okot adó eredmények.

Kompozit kockázati index

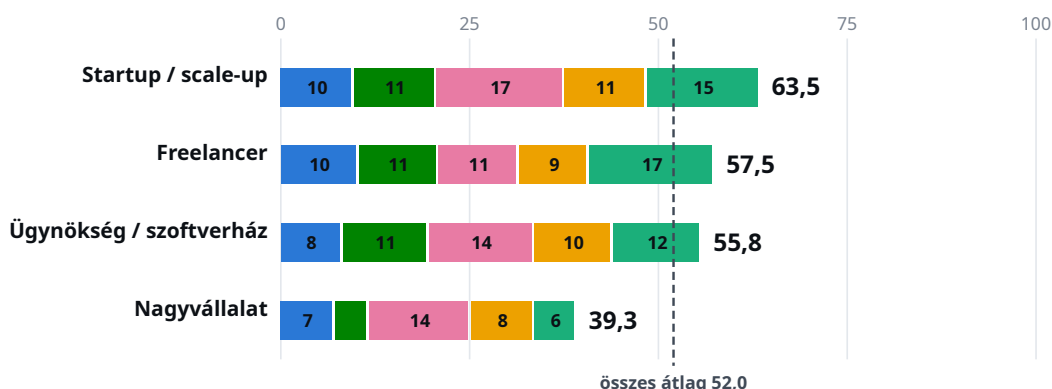
● ÍGY OLVASD – KOCKÁZATI INDEX

Egyetlen **0-100-as pontszám** öt tényezéből, egyenlő súllyal: shadow AI, hiányzó szabályzat, kizárólag publikus felhő, gyenge LLM-biztonsági felkészültség, képzéshiány. **Magasabb = kockázatosabb.** A színes szeletek megmutatják, melyik tényező hajtja a pontszámot. A szám nem abszolút mérce, hanem a cégtípusok összevetésére való.

AI-kockázati index cégtípusonként

0–100 · öt faktor egyenlő súllyal

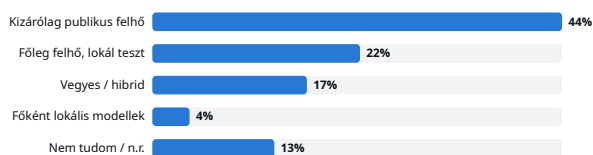
Shadow-AI szivárgás Policy-hiány Publikus-felhő kitettség LLM-biztonsági rés Képzéshiány



A sorrend megerősíti a governance-fejezet képét: a legmagasabb kockázat a **startup / scale-up** szegmensben (63,5/100), a legalacsonyabb a **nagyvállalat**nál (39,3). Nem a technológiai lemaradás termeli a kockázatot, hanem a szervezeti keret hiánya.

Lokális vs. felhő LLM-ek

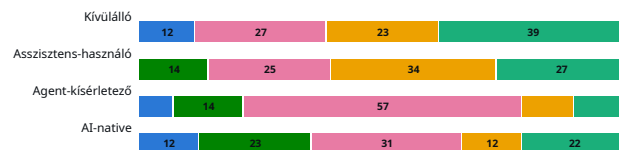
q10 · adatvédelmi megfontolás



Biztonsági érettség archetípusonként

soronként 100% · archetípus × q14

Red teaming a release előtt Beépített védelem
Alap tudatosság, nincs teszt Nem foglalkoznak vele
Nem fejleszt LLM-integrációt



09 – TANULÁSI IGÉNYEK

Mit akarnak tanulni — és mi hiányzik valójában

A kérdés az volt, hol fejlesztenék leginkább az AI-tudásukat. A válaszok szintenként élesen elválnak: a tanulási igény együtt érik a használat mélységével.

Hol fejlesztené az AI-tudását

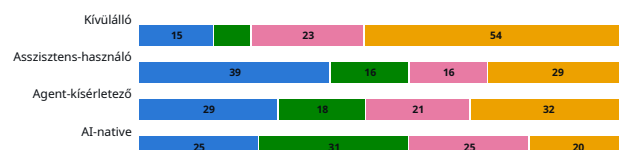
q15 · a teljes minta



Tanulási igény archetípusonként

soronként 100% · archetípus × q15

Prompt engineering AI-alapú architektúra
LLM-integráció termékbe AI biztonság & etika



A létra tisztán kirajzolódik. A **kívülálló** 54%-a az AI biztonságát és etikáját választja — kívülről nézve ez a téma. Az **asszisztens-használók** 39%-a a prompt engineeringet: még a hatékonyabb használat a cél. Az **AI-native-ok** viszont már 31%-ban az **AI-alapú architektúrát** jelölik, plusz 25% az LLM-integrációt — náluk a kérdés már rendszerszintű.

Van egy feltűnő hiány is. Az AI-native csoportban a legalacsonyabb az **AI biztonság & etika** iránti érdeklődés (20%) — miközben épp ez a csoport a legkitettebb: náluk a leggyakoribb a rendszeres shadow AI, és 43%-uknak nincs LLM-biztonsági tesztje. A kereslet és a kockázat itt ellentétes irányba mutat.

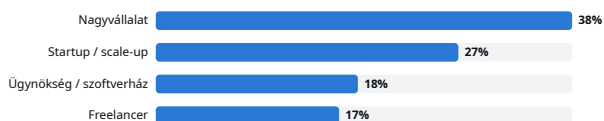
10 — MÓDSZERTAN & A MINTA

Mit mér ez a riport, és mit nem

Önkéntes kitöltésű, önbevallásos online kérdőív. Nem valószínűségi minta: **nem reprezentatív** a magyar fejlesztői populációra, és nem is annak szánjuk.

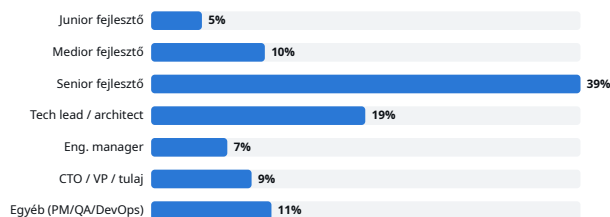
Cégtípus

q12 · a válaszadók megoszlása



Szerepkör

q12b · a válaszadók megoszlása



● AMIT AZ ADAT NEM BÍR EL

Az arányok szintje felfelé torzít. A kitöltők önként jelentkeztek egy AI-témájú kérdőívre, és a minta a senior, döntéshozói rétegre húz — ezért az abszolút százalékok felső becslésnek tekintendők, nem a teljes fejlesztői populáció értékeinek.

Minden mutató önbevallás. Az „AI-generált kódarány”, a „review-teher” és a shadow AI gyakorisága azt méri, amit a válaszadó *gondol* vagy *vállal*. A shadow AI jellemzően alulbevallott.

Az együttjárás nem ok-okozat. A korrelációk irányt mutatnak; keresztmetszeti adatból nem állapítható meg, mi mit okoz.

A kis cellák zajosak. A kis létszámú szegmenseknél a százalékok néhány fős elmozdulásra is érzékenyek — ezeknél a trend olvasható, a pontos érték nem.

Mi mozog együtt mivel

spearman-rangkorreláció · -1 (ellentétes) ... 0 (nincs) ... +1 (együtt mozog)

	Használati gyak.	Agent-autonómia	AI-kódarány	Eszközök száma	Céges támogatás	Governance-index	Review-teher
Használati gyak.	1	,44	,50	,41	,38	,16	-,14
Agent-autonómia	,44	1	,67	,24	,40	,15	-,16
AI-kódarány	,50	,67	1	,13	,36	,02	-,26
Eszközök száma	,41	,24	,13	1	,20	,28	,09
Céges támogatás	,38	,40	,36	,20	1	,36	-,02
Governance-index	,16	,15	,02	,28	,36	1	,11
Review-teher	-,14	-,16	-,26	,09	-,02	,11	1

A **kék** együtt mozgást jelent, a **piros** ellentéteset, a halvány cella azt, hogy nincs érdemi kapcsolat. A 0,2 alatti abszolút értékek gyakorlatilag kapcsolat-hiányt jelentenek.

A mátrix két dolgot mond ki. Egy: van egy koherens „**mélységi**” **tengely** — használati gyakoriság, agent-autonómia és AI-kódarány együtt mozognak (a legerősebb kapocs $p=0,67$). Kettő: a **governance-index ehhez a tengelyhez alig kapcsolódik** (AI-kódarány $\times$ governance: $p=0,02$), a review-teher pedig enyhén *negatív* irányban áll vele ($p=-0,26$).

KÉRDÉSEK

16 kérdés: használati gyakoriság (q01), eszközök (q02, többválaszos), agent-autonómia (q03), átvett munkafázisok (q04, max. 2), AI-kódarány (q05), review-teher (q06, 1–5), céges támogatás (q07, 1–5), AI-policy (q08), shadow AI (q09), lokális vs. felhő (q10), fő akadály (q11), cégtípus (q12), szerepkör (q12b), képzés (q13), LLM-biztonsági felkészültség (q14), tanulási igény (q15).

SZÁRMAZTATOTT MUTATÓK

Archetípus: szabályalapú besorolás q01 + q03 alapján. **AI-érettségi index:** q01 (35%) + q03 (35%) + q05 (30%). **Governance-index:** q08 (40%) + q13 (30%) + q14 (30%).

Kockázati index: q09 + q08 + q10 + q14 + q13, egyenlő súllyal, 0–100-ra skálázva. A nem értelmezhető válaszoknál a súlyok a maradék komponensekre normálódnak; a skálaátlagoknál a „nem releváns” válaszok kimaradnak.

ISMÉTELHETŐSÉG

Minden szám a nyers `all.csv`-ből, egyetlen szkripttel (`compute.py`) áll elő; a riport szövegébe ágyazott értékek is ugyanabból a JSON-ból jönnek, kézi átírás nélkül.

State of AI Dev 2026 – magyar fejlesztői körkép. Az adatok önbevallásosak; a minta önkéntes kitöltésű, nem reprezentatív a magyar fejlesztői populációra, és a tapasztaltabb, AI-affin rétegre torzít. Az abszolút arányok felső becslések; a megbízhatóbb jelzés a szegmensek közti különbség. A „nem releváns / nincs rálátás” válaszokat a skálaátlagoknál kihagytuk. Diagramok: saját inline-SVG renderelés, világos/sötét témához igazodva.