

Politika odpovědného přístupu k AI

INFORMAČNÍ TECHNOLOGIE | VYDÁNO: 21. ledna 2026 | Revidováno: NENÍ
K DISPOZICI

Tato politika stanovuje závazek společnosti Magna k zodpovědnému používání a vývoji řešení umělé inteligence v rámci našich operací, produktů a obchodních procesů. Tato řešení představují příležitosti ke zvýšení hodnoty, ale mohou také vytvářet právní, reputační a další rizika pro společnost Magna, její zaměstnance a další zainteresované skupiny. Pečlivé dodržování této politiky z vaší strany je zásadní pro podporu bezpečného a zodpovědného používání umělé inteligence s cílem maximalizovat příležitosti a minimalizovat rizika s ní spojená.

UPLATŇOVÁNÍ POLITIKY

Tato politika se vztahuje na společnost Magna International Inc. a všechny její provozní skupiny, divize (včetně řízených společných podniků), dceřiné společnosti a další subjekty po celém světě. Tato politika se vztahuje i na všechny osoby jednající jménem společnosti Magna, včetně zaměstnanců společnosti Magna na plný nebo částečný úvazek, nezávislých dodavatelů, vedoucích pracovníků, ředitelů, konzultantů a zástupců – v této politice označujeme všechny tyto osoby jako „vy“ nebo „pracovníci Magna“.

Tato politika stanoví požadavky na zodpovědné, etické a bezpečné používání a vývoj:

- všech **řešení AI** včetně **tradičních systémů AI**, **systémů generativní AI**, **AI agentů** a **agentické AI** a
- veškerých **výstupů AI** generované **řešením AI**.

„**Používání**“ **řešení AI** znamená, že používáte stávající **řešení AI** vyvinuté společností Magna nebo třetí stranou k provádění:

- jednoduchých úkolů, jako je generování e-mailu nebo obrázku, shrnutí textu nebo odpověď na otázku; nebo
- složitějších úkolů, jako je podpora vývoje produktů Magna, obchodních procesů a iniciativ automatizace výroby.

„**Vývoj**“ **řešení AI** znamená, že máte s řešením vyšší úroveň interakce, což zahrnuje následující body:

- vývoj specifikací pro **řešení AI**;
- sestavování, návrh, zdokonalování nebo trénování modelu a algoritmu **řešení AI**;
- shromažďování, zpracovávání nebo trénování datových sad, které jsou základem **řešení AI**;
- testování nebo ověřování **řešení AI**;
- využívání nástroje třetí strany pro **vývoj**.

Vývoj řešení AI může být:

- výhradně pro interní účely (např. automatizace výroby);
- pro použití externími zainteresovanými skupinami (např. zákazníky, dodavateli nebo jinými obchodními partnery); nebo
- k integraci do produktů Magna.

Vývoj řešení AI může být navíc prováděn:

- interně pracovníky Magna (např. prostřednictvím nástrojů třetích stran, jako je NVIDIA Isaac Sim, a také

prostřednictvím funkcí **AI** v tradičních systémech, jako je CATIA);

- externím partnerem pro společnost Magna nebo jejím jménem; nebo
- společně týmem, který zahrnuje pracovníky Magna a personál externího vývojáře.

Kromě výrazů „používání“ a „vývoj“ jsou v této politice tučně zvýrazněny i další důležité pojmy. Definice těchto důležitých pojmů naleznete v **Dodatek A** k této politice.

ODPOVĚDNÝ PŘÍSTUP K AI VE SPOLEČNOSTI MAGNA

Dosažení přínosů produktivity díky **AI** a zároveň minimalizace rizik a ochrana našich zaměstnanců jsou pro společnost Magna důležité. Vaše role je důležitá, aby společnost Magna mohla dosahovat těchto cílů. K tomu je třeba pečlivě dodržovat níže uvedené zásady a postupy:

1. Pochopit a minimalizovat rizika spojená s používáním řešení AI

Používání **řešení AI** může vytvářet rizika, která je třeba pochopit a snažit se je minimalizovat. Jedná se například o tato:

- **Zkreslení:** **Řešení AI** se mohou učit a posilovat zkreslení přítomná v jejich trénovacích datech, což může vést ke zkresleným výsledkům a nezamýšleným důsledkům, jako jsou diskriminační praktiky při náboru. Rizika zkreslení je třeba si uvědomovat a při posuzování **výstupů AI** postupovat kriticky.
- **Kybernetické hrozby:** **Řešení AI** lze použít nebo zneužít k zahájení kybernetických útoků, krádeží identit a ohrožení bezpečnosti. **Řešení AI** nikdy nepoužívejte k těmto ani jiným nezákonným účelům. Kromě toho pomůžete chránit společnost Magna před kybernetickými útoky a dalšími externími hrozbami tím, že pro **řešení AI** nebudete instalovat neschválené pluginy, konektory, doplňky ani API.
- **Nedostatečná transparentnost:** Algoritmy **AI** fungují způsobem, který není dobře chápán, a často produkují chybné **výstupy AI**. Jak je podrobněji uvedeno níže, nad **výstupy AI** je nutné dohlížet, abyste zajistili, že jsou přesné, spolehlivé, relevantní, nezkrácené a vhodné pro zamýšlené použití.
- **Důvěrnost a ochrana osobních údajů:** Mnoho **řešení AI** používaných údaje shromažďuje, což vývojářům potenciálně umožňuje přístup k důvěrným a/nebo právně chráněným datům a osobním informacím. Pro účely trénování, vývoje nebo revizí modelů je nutné odmítnout souhlas s použitím **vstupů vlastních nebo kontrolovaných společností Magna** (včetně osobních údajů pracovníků Magna) **systémem generativní AI** třetí strany.
- **Práva duševního vlastnictví:** **Systémy generativní AI** mohou vytvářet **výstupy AI**, které porušují práva duševního vlastnictví třetích stran. Mezi příklady patří použití psaného textu bez řádného uvedení zdroje a citace, neoprávněné použití nebo kopírování obrázků, uměleckých děl, návrhů, softwarového kódu nebo jiných materiálů a návrhy produktů nebo procesů, které porušují patenty třetích stran.

Řešení AI a výstupy AI nesmíte používat způsobem, který porušuje práva duševního vlastnictví třetích stran (včetně autorských práv, patentů, návrhů a ochranných známek).

Pokud navíc používáte **řešení AI** k vytvoření jakéhokoli pracovního produktu, u kterého může Magna vyžadovat ochranu podle zákonů o duševním vlastnictví, musíte spolupracovat s právním týmem Magna pro duševní vlastnictví, abyste zajistili, že Magna bude moci chránit jakýkoli pracovní produkt, který může být založen na **výstupech AI**.

2. Věnujte se dohledu

Výstupy AI vyžadují kritickou lidskou kontrolu, než budou finalizovány, než se na ně kdokoli bude spoléhat nebo než budou interně či externě sdíleny. Protože jste zodpovědní za veškeré **výstupy AI**, které používáte ve své práci pro společnost Magna, je nutné použít promyšlené vyhodnocení, analytické myšlení a rozumný úsudek, než se spolehnete na **výstupy AI**.

Nespoléhejte se na **výstupy AI**, o kterých víte, že byly vygenerovány při porušení platných právních předpisů nebo práv kohokoli jiného, případně že by mohly poškodit pověst společnosti Magna. Příklady mohou zahrnovat **výstupy AI**, které porušují autorská práva nebo práva k ochranným známkám jiných osob.

Pokud si i po vynaložení přiměřeného úsilí stále nejste jisti platností **výstupu AI**, nespolehejte se na něj – zejména v situacích, které představují významné právní nebo reputační riziko.

3. Používejte řešení AI schválená společností Magna

Společnost Magna ověřila řadu **řešení AI**, aby bylo zajištěno dodržování našich standardů kybernetické bezpečnosti a našich politik ochrany důvěrnosti a ochrany osobních údajů. Kdykoli je to možné, měli byste používat **řešení AI schválená společností Magna**, jako jsou tato:

- **MAVIS**;
- Asistenti Microsoft M365 Copilot a Copilot Chat;
- Asistent kódování a programování GitHub Copilot;
- Platforma Microsoft Azure AI Foundry Models;
- Služba Amazon Bedrock pro přístup k základním modelům a jejich správu;
- Platforma Databricks.

Pokud chcete používat **řešení AI** a máte v úmyslu je používat se **vstupy vlastněnými nebo kontrolovanými** společností Magna, je nutné nejprve ověřit, zda je uvedeno jako **schválené řešení AI** na **AI MagNET – řešení a technologie AI**. Pokud již nebyl schválen, jste zodpovědní za navržení nástroje na platformě **Magnet** v souladu s příslušným schvalovacím procesem.

Vždy předpokládejte, že **vstupy** zadané do neschváleného **řešení AI** nebudou nikdy zpracovávány jako důvěrné. V důsledku toho by použití **vstupů** v neschváleném **řešení AI** mohlo ohrozit práva duševního vlastnictví společnosti Magna.

4. Zakázané použití

Nesmíte používat **řešení AI** k jakémukoli jednání nebo dosažení jakéhokoli výsledku, který je zakázán zákonem nebo který by poškodil pověst společnosti Magna, a to včetně jakékoli formy podvodu, diskriminace, obtěžování, zastrasování, šikany nebo jiné újmy. Dále nesmíte používat **řešení AI** k žádnému z následujících účelů:

- manipulace s chováním nebo zneužití zranitelnosti jakékoli osoby škodlivým způsobem, včetně zranitelností souvisejících s věkem, postižením nebo specifickou sociální či ekonomickou situací;
- hodnocení nebo klasifikace osoby způsobem, který by mohl vést k diskriminaci nebo újmě, a to i na základě sociálního chování nebo osobních charakteristik;
- vytváření databází pro rozpoznávání obličejů (emocí) prostřednictvím necíleného sběru snímků obličejů;
- rozpoznávání emocí osob na pracovišti; nebo
- kategorizace jakékoli osoby na základě citlivých nebo zákonem chráněných osobních charakteristik, včetně etnického původu, náboženství, věku, sexuální orientace a politických názorů, způsobem, který by mohl způsobit diskriminaci nebo újmu.

5. Zahrnutí prohlášení o vyloučení odpovědnosti za AI

V každé situaci týkající se **řešení AI**, kdy by nedodržení tohoto ustanovení mohlo reálně způsobit škodu nebo uvést jiné v omyl.

- **Zákonný požadavek:** V jakékoli situaci, kdy **platné právní předpisy** vyžadují zahrnutí prohlášení o vyloučení odpovědnosti, **musíte** tak učinit.
- **Osvědčený postup:** Pro běžné položky, jako jsou každodenní e-maily, je prohlášení o vyloučení odpovědnosti obecně zbytečné. V jiných veřejně exponovaných nebo pro podnikání kritických situacích, jako je práce z velké části vygenerovaná **AI**, a/nebo která by mohla významně ovlivnit rozhodování, **je třeba** prohlášení o vyloučení odpovědnosti.

Ostatní specifické okolnosti, za kterých **musíte** prohlášení o vyloučení odpovědnosti za **AI** použít, například tyto:

- Pokud **řešení AI** interagují přímo s lidmi (např. chatboti) a přiměřeně informované osobě není zřejmé, že interakce zahrnuje **řešení AI**;
- Pokud **výstupy AI**, které používáte nebo odesíláte, zahrnují upravené obrázky, zvuk nebo video, které vypadají jako autentické nebo by mohly přiměřeně vnímané jako autentické, přestože nejsou;
- Pokud jsou **výstupy AI** určené k informování interních nebo externích zainteresovaných skupin o záležitostech veřejného zájmu (např. sdělení o produktech, službách nebo bezpečnosti, která mohou mít vliv na zákazníky nebo jiné zainteresované skupiny) prezentovány jako věcné nebo směrodatné.

Doporučená forma prohlášení o vyloučení odpovědnosti: „Následující obsah byl vygenerován pomocí řešení AI.“

Platí určité výjimky. Například text generovaný **AI** nevyžaduje označení, pokud byl zkontrolován a upraven člověkem, který přebírá redakční odpovědnost. Navíc není nutné žádné prohlášení o vyloučení odpovědnosti, pokud je již zřejmé, že interakce zahrnuje **AI**.

6. Akvizice nebo vývoj řešení AI

Pokud máte v úmyslu pořídit si (včetně zakoupení, získání licence nebo předplacení) nebo vyvinout/spoluvyvinout **řešení AI**, jste zodpovědní za dodržení schvalovacího procesu společnosti Magna pro **řešení AI**, který je k dispozici prostřednictvím odkazu na platformě **AI MagNET – řešení a technologie AI**.

Tento proces zahrnuje dva oddělené kroky:

- **Krok 1:** sledování plánovaných iniciativ, nápadů a případů užití v oblasti **AI** za účelem podpory transparentnosti, posouzení obchodní hodnoty a finanční proveditelnosti, minimalizace duplicity úsilí a posílení správy a řízení rizik;
- **Krok 2:** dokumentace, kategorizace a posouzení rizik **řešení AI** z předchozího kroku, které jsou schváleny a implementovány.

V závislosti na vaší roli ve vztahu k pořizovanému nebo vyvíjenému **řešení AI** můžete být považováni za „vývojáře“ řešení, a to i pro účely dodržování této politiky a **zásad odpovědného přístupu k AI** společnosti Magna, stejně jako **postupů pro vývojáře AI** a **EUAIA**.

7. Agenti AI

Můžete používat **agenty AI**, kteří byli vytvořeni pro vaše specifické pracovní potřeby prostřednictvím **schválených řešení AI**. **Agenti AI** jsou **řešení AI**, která plně podléhají této politice.

Pokud jste majitelem firmy a nasazujete **agenta AI** pro své vlastní použití nebo pro použití svého týmu, musíte si tohoto **agenta AI** nechat nejprve schválit prostřednictvím dvoustupňového schvalovacího procesu uvedeného výše v části 6. Dále musíte spolupracovat s technickým týmem při požadovaném monitorování a dohledu nad vašimi **agenty AI** po celou dobu jejich životního cyklu (tj. od počátečního návrhu až po vyřazení, včetně komunikace mezi agenty), což zahrnuje dodržování požadavků na správu a řízení, organizačních IT standardů a pokynů, i **zásad RAI**. Nakonec budete muset splnit i všechny další požadavky, které uvádějí **postupy pro vývojáře AI** ve vztahu k **agentům AI**.

DALŠÍ INFORMACE

1. Dodržování politik společnosti Magna a platných zákonů

Ve všech situacích týkajících se **používání** nebo **vývoje řešení AI** musíte dodržovat **platné právní předpisy**, **zásady odpovědného přístupu k AI** společnosti Magna a další relevantní politiky společnosti Magna (např. politiku ochrany důvěrných informací, politiku ochrany osobních údajů atd.). Další informace uvádí **dodatek B** této politiky, kde naleznete seznam takových zákonů a souvisejících politik společnosti Magna.

Kromě toho se povinnosti společnosti Magna a každé osoby, na kterou se tato politika vztahuje, budou řídit **platnými právními předpisy** a v případě jakýchkoli nesrovnalostí budou vykládány co nejbližší požadavkům této politiky a zároveň zůstanou v souladu s **platnými právními předpisy**.

2. Dodržování požadavků zákazníků, dodavatelů, prodejců a dalších třetích stran a dohod o mlčenlivosti

Při nakládání s daty a informacemi třetích stran v rámci **řešení AI** musíte zvážit smluvní omezení, restrikce a zákazy. V některých případech nemusí být povoleno používat a/nebo vkládat taková data třetích stran do **řešení AI**, protože by to mohlo být v rozporu s dohodou o mlčenlivosti nebo jinými ujednáními o důvěrnosti. S právním zástupcem skupiny, regionu nebo společnosti byste měli prodiskutovat příslušné povinnosti týkající se mlčenlivosti, omezení, zákazů a dalších smluvních omezení, abyste společnosti Magna pomohli vyhnout se jakémukoli porušení smluvních závazků. Kromě toho musíte dodržovat také politiku důvěrnosti společnosti Magna.

3. Školení

Magna vám poskytne zdroje a školení, které vám pomohou lépe porozumět možnostem a omezením **řešení AI**. Jste zodpovědní za včasné absolvování veškerých požadovaných školení.

4. Reporting

Monitorujte, dokumentujte a hlase jakékoli poruchy, se kterými se při vývoji **řešení AI** setkáte, a také jakékoli zkreslené, škodlivé, nepřesné nebo anomální **výstupy AI**. V mnoha případech může k takovému reportování docházet přímo v rámci **řešení AI** – v takovém případě funkci nahlášení použijte. V situacích, kdy v samotném **řešení AI** neexistuje funkce pro hlášení chyb nebo se domníváte, že rizika/důsledky jsou významné, můžete chybu nahlásit na **[horké lince Magna](#)** v části **Obavy týkající se ochrany osobních údajů / umělé inteligence (AI)**.

Pokud se dozvíte o jakémkoli podezření na porušení této politiky nebo o porušení práv duševního vlastnictví společnosti Magna, je nutné to nahlásit prostřednictvím **[horké linky Magna](#)**.

V závislosti na regionu a okolnostech může být nutné podat hlášení příslušným vládním orgánům; žádné ustanovení této politiky vám nebrání v tom, abyste s vládními orgány hovořili sami, individuálně.

5. Monitorování dodržování předpisů

Společnost Magna si vyhrazuje právo sledovat dodržování této politiky ze strany pracovníků Magna. Magna zavede opatření k prevenci, monitorování a reakci na incidenty zahrnující **řešení AI**, včetně případů újmy a zkreslených výstupů, jakož i narušení bezpečnosti, důvěrnosti a ochrany osobních údajů.

6. Odpovědnost za porušení této politiky

Pokud podmínky těchto zásad porušíte, hrozí vám disciplinární postih, kterým může být až ukončení pracovního poměru.

7. Další informace

O další informace nebo pokyny se prosím obraťte na své určené **[globální lídry v oblasti AI](#)** nebo napište na adresu **ai.governance@magna.com**.

Datum vydání:	21. ledna 2026
Revize proběhla:	n/a
Příští revize:	1. čtvrtletí 2027
Vydáno:	Personální oddělení a informační technologie
Schváleno:	Personální oddělení a informační technologie

PŘÍLOHA A – DEFINICE

„**Agentická AI**“ označuje **řešení AI**, které dokáže plánovat, provádět akce nebo plnit úkoly s určitou autonomií, a přitom stále funguje pod lidským dohledem (tj. **agentická AI** je schopnost (AI, která může jednat, plánovat nebo sledovat cíle)).

„**AI**“ znamená umělá inteligence.

Jako „**agenti AI**“ se označují specifické nástroje nebo systémy, které využívají agentické schopnosti k provádění úkolů nebo pracovních postupů jménem uživatele v rámci kontrolovaných parametrů (tj. **agenti AI** jsou nástroje nebo systémy, které tuto schopnost využívají k provádění úkolů v praxi). Příklady: Agenti **AI** v Microsoft Teams, kteří pomáhají s plánováním schůzek nebo zodpovídáním otázek zaměstnanců; specializovaní asistenti **AI** atd.

„**Postupy pro vývojáře umělé inteligence**“ označují **Postupy pro vývojáře umělé inteligence**.

„**Výstupy AI**“ se rozumí jakýkoli obsah (včetně textu, obrázků, zvuku, videa a softwarového kódu), výsledky, doporučení, rozhodnutí a/nebo jiné výstupy generované **řešením AI**.

„**Řešení AI**“ označuje všechny systémy, funkce, případy užití, produkty, platformy a nástroje AI, včetně **tradičních systémů AI**, **systémů generativní AI**, **agentů AI** a **agentické AI**.

„**Platné právní předpisy**“ definuje **dodatek B**.

„**Schválená řešení AI**“ znamenají **řešení AI**, která odpovídají bezpečnostním protokolům IT společnosti Magna, splňují požadavky na důvěrnost a byla schválena a nasazena pro použití v podnikání společnosti Magna. Seznam **schválených řešení AI** společnosti Magna naleznete **[zde](#)**.

„**Vývoj řešení AI**“ znamená, že máte s řešením vyšší úroveň interakce, což zahrnuje následující body:

- vývoj specifikací pro **řešení AI**;
- sestavování, návrh, zdokonalování nebo trénování modelu a algoritmu **řešení AI**;
- shromažďování, zpracovávání nebo trénování datových sad, které jsou základem **řešení AI**;
- testování nebo ověřování **řešení AI**;
- využívání nástroje třetí strany pro **vývoj**.

Vývoj řešení AI může být:

- výhradně pro interní účely (např. automatizace výroby);
- pro použití externími zainteresovanými skupinami (např. zákazníci, dodavatelé nebo jinými obchodními partnery); nebo
- k integraci do produktů Magna.

Vývoj řešení AI může být navíc prováděn:

- interně pracovníky Magna (např. prostřednictvím nástrojů třetích stran, jako je NVIDIA Isaac Sim, a také prostřednictvím funkcí **AI** v tradičních systémech, jako je CATIA);
- externím partnerem pro společnost Magna nebo jejím jménem; nebo
- společně týmem, který zahrnuje pracovníky Magna a personál externího vývojáře.

„**EUAIA**“ znamená *zákon EU o AI*.

„**Systémy generativní AI**“ označují **řešení AI**, která jsou primárně navržena ke generování nového obsahu (např. textu, obrázků, zvuku, videa nebo kódu atd.) nebo k autonomnímu plánování a provádění více krokových akcí směrem k cílům na základě vzorců získaných z dat, spíše než k pouhé analýze nebo klasifikaci stávajících dat.

„**Vstupy**“ znamenají jakékoli vstupy, data, dotazy, příkazy, informace nebo dokumenty zadané do **řešení AI**.

„**Magna**“ označuje společnost Magna International Inc. a všechny její provozní skupiny, divize (včetně řízených společných podniků), dceřiné společnosti a další subjekty po celém světě.

„**Vstupy vlastněné nebo kontrolované společnostmi Magna**“ znamenají data, informace nebo dokumenty, které jsou v péči, držení nebo pod kontrolou společnosti Magna a patří společnosti Magna, pracovníkům Magna nebo zákazníkům, dodavatelům či jiným obchodním partnerům společnosti Magna.

Jako „**pracovníci Magna**“ nebo „**vy**“ se označují všechny osoby, které jednají jménem společnosti Magna, včetně zaměstnanců společnosti Magna na plný nebo částečný úvazek, nezávislých dodavatelů, vedoucích pracovníků, ředitelů, konzultantů a zástupců.

„**MAVIS**“ označuje asistenta virtuálního informačního systému AI Magna.

„**Zásady RAI**“ označují **zásady odpovědného přístupu k AI**.

„**Tradiční systémy AI**“ označují **řešení AI**, která využívají statistické/strojové učení a/nebo logiku založenou na pravidlech k analýze nebo klasifikaci stávajících dat a vytváření výstupů, jako jsou předpovědi, doporučení nebo rozhodnutí, a nejsou primárně navržena ke generování nového obsahu ani k autonomnímu plánování a provádění více krokových akcí směrem k cílům.

„**Používání**“ **řešení AI** znamená, že používáte stávající **řešení AI** vyvinuté společností Magna nebo třetí stranou k provádění:

- jednoduchých úkolů, jako je generování e-mailu nebo obrázku, shrnutí textu nebo odpověď na otázku; nebo
- složitějších úkolů, jako je podpora vývoje produktů Magna a iniciativ automatizace výroby.

PŘÍLOHA B – PLATNÉ PRÁVNÍ PŘEDPISY A POLITIKY SPOLEČNOSTI MAGNA

Platné právní předpisy

„**Platné právní předpisy**“ označují veškeré zákony, stanovy, předpisy, pravidla, vyhlášky, kodexy, směrnice, nařízení, rozsudky, pokyny a vládní požadavky, které mohou příležitostně nabýt účinnosti a které jsou relevantní pro výběr, používání, vývoj nebo nasazování **řešení AI**, včetně těchto:

- Nařízení Evropského parlamentu a Rady (EU) 2024/1689 ze dne 13. června 2024, kterým se stanoví harmonizovaná pravidla pro umělou inteligenci („EUAIA“)

Relevantní/související politiky společnosti Magna

- **Politika používání důvěrných informací Magna**
- **Magna International Inc. Politika zveřejňování informací**
- **Magna International Inc. Politika, postupy a pokyny k důvěrnosti a ochraně osobních údajů**
- **Politika klasifikace informací Magna**
- **Politika zabezpečení společnosti Magna**
- **Globální politika společnosti Magna pro e-mail, internet/intranet a sociální média**
- **Politika zabezpečení IT/OT společnosti Magna**
- **Magna International Inc. Politika reakce na incident kyberbezpečnosti**
- **Politika ochrany zdraví, bezpečnosti a životního prostředí společnosti Magna**