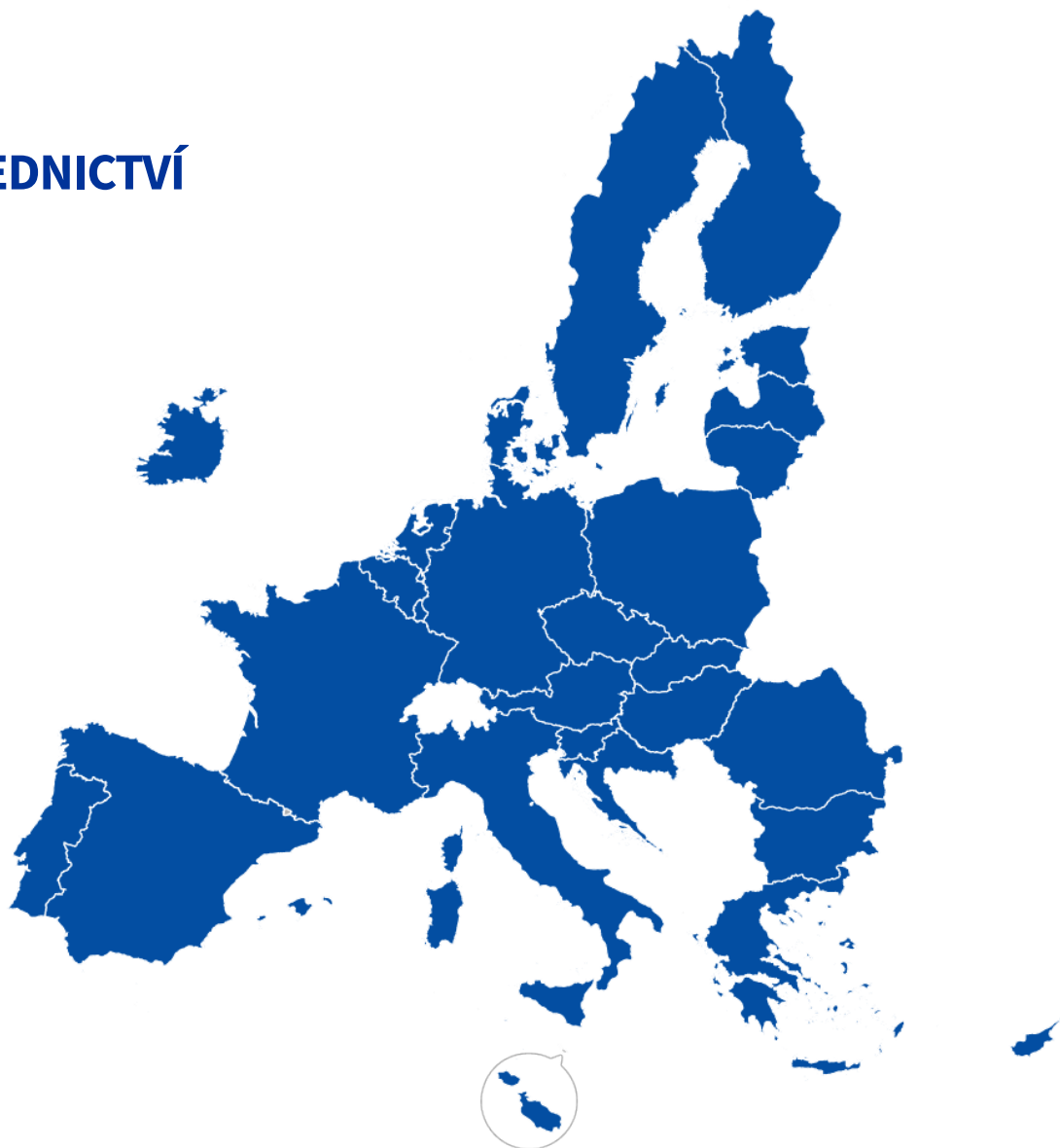




PŘEDSEDNICTVÍ



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Dohodnout se na společném postoji Rady dvojí většinou (tzn. alespoň 55 % členských států zastupujících nejméně 65 % celkového počtu obyvatel EU). V ideálním případě je cílem dosažení jednomyslného rozhodnutí; pravidlo „dvojí většiny“ se uplatní v případě, že jednomyslného rozhodnutí dosáhnout nelze. V praxi se o většině rozhodnutí nehlasuje – předsednictví se snaží nalézt kompromis, který získá podporu všech delegací.



Vaše úloha

Připojte se k delegaci členského státu, který aktuálně vykonává předsednictví. Lektor vám sdělí, o kterou zemi se jedná.

V rámci delegace máte za úkol předsedat oficiálním jednáním, zatímco vaši kolegové prezentují zájmy a postoje dané země. Dohlížíte na všechna oficiální jednání, udělujete slovo delegátům a zajišťujete dodržování časového harmonogramu.

Zahájení

1. **Zasedání byste měli formálně zahájit**, přivítat na něm všechny delegáty a Komisi a vysvětlit jim jeho účel (maximálně 1 minuta). Může vám pomoci text uvedený v šedém rámečku na následující straně.
2. Poté **požádáte delegáta Komise, aby návrh prezentoval**.
3. Budou **následovat (maximálně 1minutová) úvodní prohlášení členských států**, jimž budete předávat slovo podle pořadí, v němž se přihlásily. Vyzvěte členské státy, aby jasně uvedly, ve kterých aspektech návrhu se jejich zájmy shodují se zájmy ostatních členských států.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Jednání

1. **Moderujte následná oficiální jednání.** Ministři mohou vystoupit pouze poté, co jim dáte slovo. Předávejte delegátům slovo v pořadí, v jakém se o něj přihlásili.
2. Požádejte delegáty, aby neopakovali argumenty a postoje, které již byly vyjádřeny. Vyzvěte členské státy, které zastávají podobné postoje, aby vždy jeden hovořil i za ostatní. **Čas je omezený!**
3. **Oficiální jednání jsou přerušována, aby mohly probíhat neformální rozhovory.** Oznámíte začátek neformálních rozhovorů a čas, kdy bude opět zahájeno formální jednání. Na neformální rozhovory nemusíte dohlížet.

Dosažení dohody

1. Během neformálních rozhovorů máte čas na **přípravu kompromisního návrhu**. Můžete se samozřejmě také zeptat členských států, zda mají ke kompromisu nějaké podněty.
2. O nejslibnějším návrhu či návrzích znění článků 1 a 2 dejte hlasovat. Společný postoj Rady se přijme dvojí většinou (55 % členských států zastupujících nejméně 65 % celkového počtu obyvatel EU). K výpočtu hlasů **použijte hlasovací kalkulačtor**. Lektor vám v této fázi pomůže.
3. Při závěrečném hlasování o kompromisním návrhu členským státům připomeňte, že v této fázi se musí rozhodnout, zda upřednostňují kompromis, se kterým možná nebudou stoprocentně spokojeny, nebo zda dávají přednost tomu, aby jednání skončilo bez výsledku.

Dámy a pánové, vážení členové Rady, paní komisařko / pane komisaři,

mám tu čest zahájit dnešní zasedání Rady.

Jak víte, _____ v současné době předsedá Radě, a bude jednání řídit. Mým úkolem bude dnešní jednání moderovat, nikoliv prezentovat stanovisko naší země k tomuto tématu.

Dnes je na pořadu jednání nové nařízení o využívání a vývoji umělé inteligence v EU. Jedná se o téma, které je pro všechny evropské občany i pro úlohu EU ve světě velmi důležité, protože _____.

Ještě než předám slovo Komisi, chtěl/a bych zdůraznit, že návrhem se zabývala právní služba Rady, která nás informovala, že je návrh v souladu se Smlouvami a má nezbytný právní základ.

Nyní si dovolíme požádat Komisi, aby prezentovala svůj návrh. Poté bude mít každý z vás možnost přednést krátké úvodní prohlášení a představit v něm postoj své země. Těšíme se na konstruktivní a plodnou diskusi.

Paní komisařko / pane komisaři, máte slovo.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

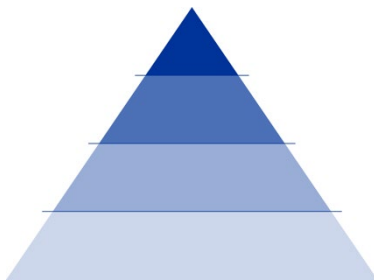
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko	na základě účelu	předběžná opatrnost		národní bezpečnost	
Belgie	na základě schopností	předběžná opatrnost		otevřený zdroj	
Bulharsko	na základě účelu	flexibilní		bez výjimky	
Chorvatsko	na základě účelu	předběžná opatrnost		startupy	
Kypr	na základě schopností	předběžná opatrnost	přezkum za 2 roky	bez výjimky	
Česko	na základě schopností	flexibilní	přezkum za 5 let	bez výjimky	
Dánsko	na základě účelu	předběžná opatrnost	finanční podpora	armáda/obrana + národní bezpečnost	
Estonsko	na základě účelu	flexibilní	testování v reálných podmínkách	armáda/obrana + národní bezpečnost	
Finsko	na základě účelu	flexibilní	odložit hlasování	startupy	
Francie	na základě účelu	předběžná opatrnost	zahrnout oblast vzdělávání	armáda/obrana + národní bezpečnost	
Německo	na základě schopností	předběžná opatrnost		bez výjimky	
Řecko	na základě schopností	předběžná opatrnost	přezkum za 3 roky	armáda/obrana + národní bezpečnost	
Maďarsko	na základě schopností	flexibilní		armáda/obrana	
Irsko	na základě účelu	flexibilní	flexibilita při testování	startupy	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie	na základě schopností	předběžná opatrnost		startupy	
Lotyšsko	na základě schopností	flexibilní	odstranit požadavek na monitorování po uvedení na	armáda/obrana	
Litva	na základě schopností	flexibilní		bez výjimky	
Lucembursko	na základě účelu	flexibilní	přezkum za 3 roky	startupy	
Malta	na základě účelu	flexibilní	testování v reálných podmínkách	národní bezpečnost	
Nizozemsko	na základě schopností	předběžná opatrnost		národní bezpečnost	
Polsko	na základě účelu	předběžná opatrnost		armáda/obrana	
Portugalsko	na základě účelu	flexibilní	odložit hlasování	otevřený zdroj	
Rumunsko	na základě účelu	flexibilní		národní bezpečnost	
Slovensko	na základě účelu	flexibilní	finanční podpora	startupy	
Slovinsko	na základě schopností	předběžná opatrnost		startupy	
Španělsko	na základě schopností	předběžná opatrnost		armáda/obrana + národní bezpečnost	
Švédsko	na základě účelu	flexibilní		armáda/obrana + národní bezpečnost	
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

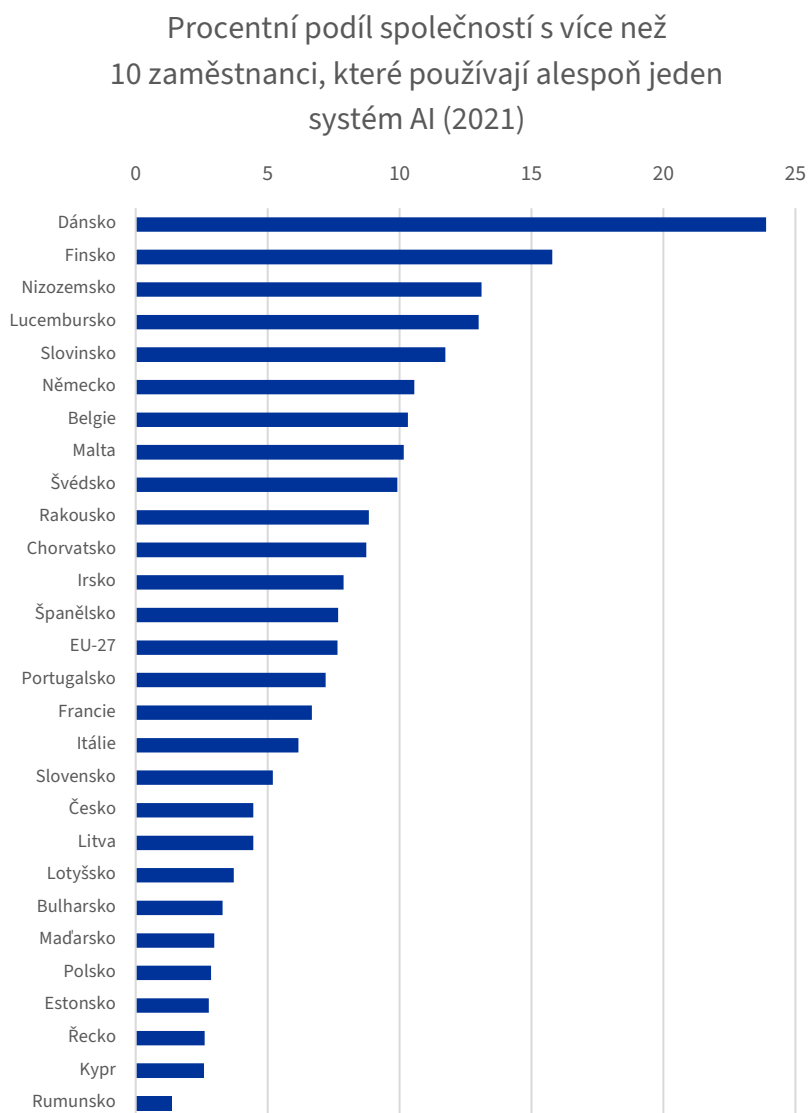
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

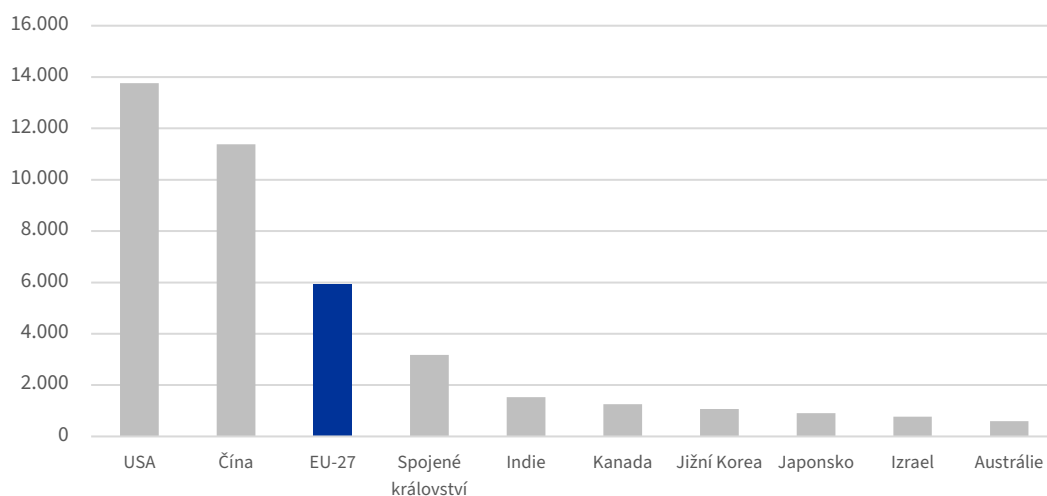
	Evropská unie
Počet obyvatel	447 695 350
HDP na obyvatele (v EUR)	38 150
Míra nezaměstnanosti	6,1 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %



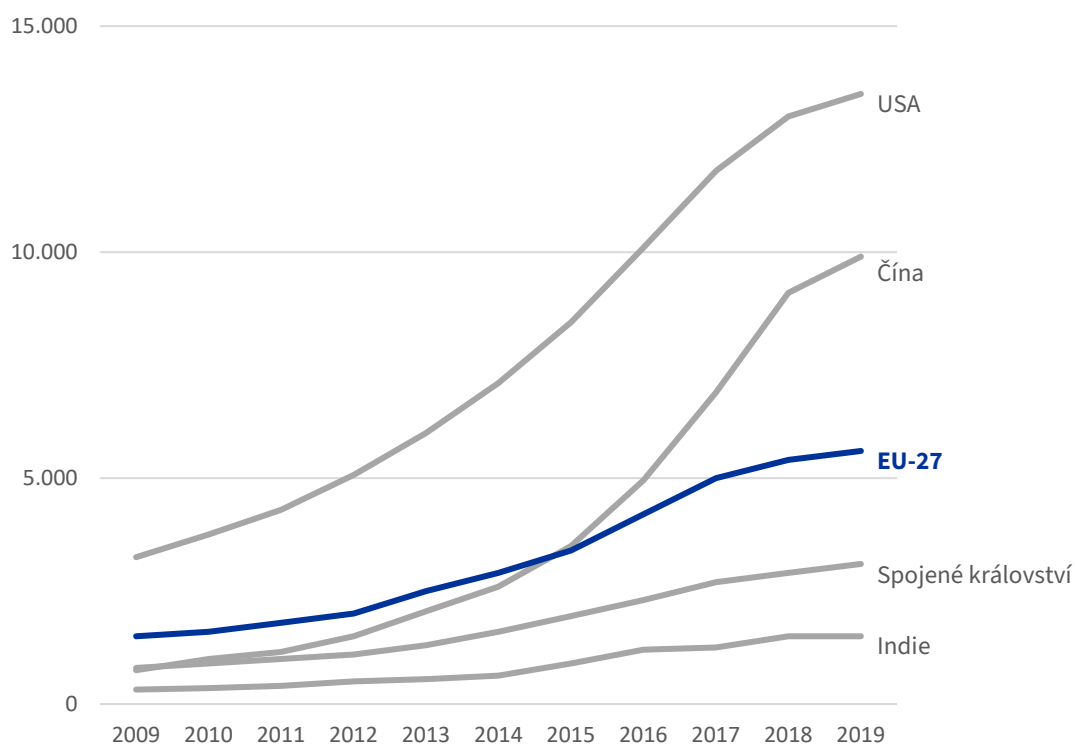
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



EVROPSKÁ KOMISE



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady váš návrh týkající se právního rámce pro využívání umělé inteligence.

Jakmile vám předsednictví (které zasedání řídí) udělí slovo, budete prezentovat svůj návrh zástupcům členských států a zapojíte se do jednání.

Váš hlavní cíl

Zajistit, aby byla do konečné verze zahrnuta co největší část původního textu, a zároveň pomoci Radě dosáhnout společného postoje. Je zásadně důležité nalézt společné evropské řešení!

Váš úkol

1. Pečlivě si **přečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Připravte úvodní** (maximálně 1minutové) **prohlášení**, v němž svůj návrh představíte a objasníte. Může vám pomoci text uvedený v šedém rámečku vpravo.
3. Pozorně vyslechněte postoje členských států. **Poznamenejte si je** pro přehlednost **do tabulky** na stranách 8 a 9.
4. Během jednání vysvětlíte svůj postoj a **svůj návrh obhajíte**. Ujměte se slova a pomozte Radě najít kompromis, který by naplňoval vaše očekávání!
5. Na konci jednání **nemáte právo hlasovat**. Během neformálních přestávek se snažte delegace přesvědčit, aby se k vašim návrhům připojily.



Vážené členové Rady, dámy a pánové,

jsme potěšeni, že se dnes uskuteční tato důležitá debata o nařízení týkajícím se umělé inteligence v EU. Toto nařízení by přineslo velké výhody občanům i Unii jako celku, protože _____

V našem návrhu nařízení navrhujeme **klasifikaci rizik** pro vysoce rizikové aplikace AI **založenou na účelu a požadavky na preventivní testování**, neboť _____

Výjimky z ustanovení článku 1 budou povoleny v případě, že systémy AI budou vyvinuty pro vojenské nebo obranné účely. A to z důvodu, že _____

Společně jako komunita můžeme pomoci zajistit, aby byla AI vyvíjena odpovědně.

Děkuji vám za pozornost.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Je třeba urychleně přijmout opatření na podporu inovací a na posílení globální konkurenceschopnosti umělé inteligence pocházející z EU a zároveň důkladně zohlednit veškerá potenciální rizika pro základní práva a ochranu spotřebitele.
- Vzhledem k tomu, že se jedná o novou technologii s nepředvídatelnými důsledky pro jednotlivce, společnost, podniky i ekonomiku, je třeba k umělé inteligenci přistupovat ambiciózně ale zároveň obezřetně. Váš návrh tento aspekt zohledňuje, neboť při definování „vysoce rizikových“ systémů AI podporuje přístup založený na účelu. Takový přístup považuje AI za vysoce rizikovou v citlivých oblastech, ve kterých se rozhoduje o životě a smrti nebo které mají významný dopad na životní podmínky občanů EU.
- Rada se již dohodla na kategoriích pro systémy AI obnášející nepřijatelné riziko, které mají být v Unii zakázány, jakož i pro systémy s nízkým a nulovým rizikem. Předmětem dnešního jednání je jasně vymezit, co představuje „vysoké riziko“. Žádný systém AI spadající do této kategorie nesmí být považován za systém s nízkým nebo nulovým rizikem.
- V rámci dnešních diskusí se také snažíte stanovit jasné požadavky na vysoce rizikové systémy AI. Rámec preventivního testování zajišťuje, že umělá inteligence je vyvíjena odpovědně. Takovéto normy zároveň vývojáře motivují k tomu, aby vytvářeli vysoce kvalitní a spolehlivé systémy uznávané po celém světě. Očekáváte, že se EU díky tomuto přístupu stane jedním z lídrů v oblasti AI a že v této oblasti rychle dožene Spojené státy a Čínu.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Zatímco prioritou je i nadále prosazování jednotného vedoucího postavení v oblasti digitálních technologií v EU a posilování rozměru zahraničního obchodu, jako je celosvětový vývoz produktů AI vyvinutých v EU, v konkrétních odvětvích, která jsou zásadní pro veřejnou bezpečnost, je nezbytné zachovat národní autonomii.
- Dochází například k tomu, že se rychle stupňují výzvy v oblasti mezinárodní bezpečnosti, kde AI hraje klíčovou úlohu mimo jiné při šíření kybernetické války a rozvoji škodlivého hackerského softwaru. Obranná protiopatření nesmějí být narušována ani těmi sebemenšími regulačními překážkami. Článek 1 by se proto neměl vztahovat na aplikace pro vojenské a obranné účely.
- Pro všechny ostatní účely nejsou zapotřebí žádné další výjimky, neboť minimální požadavky na testování zajišťují, že AI je vyvíjena odpovědně.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

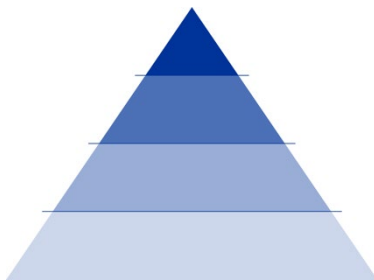
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zveřejnění výsledků testů regulačním orgánům EU.
- Monitorování po uvedení na trh za účelem sledování výkonnosti.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování.
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

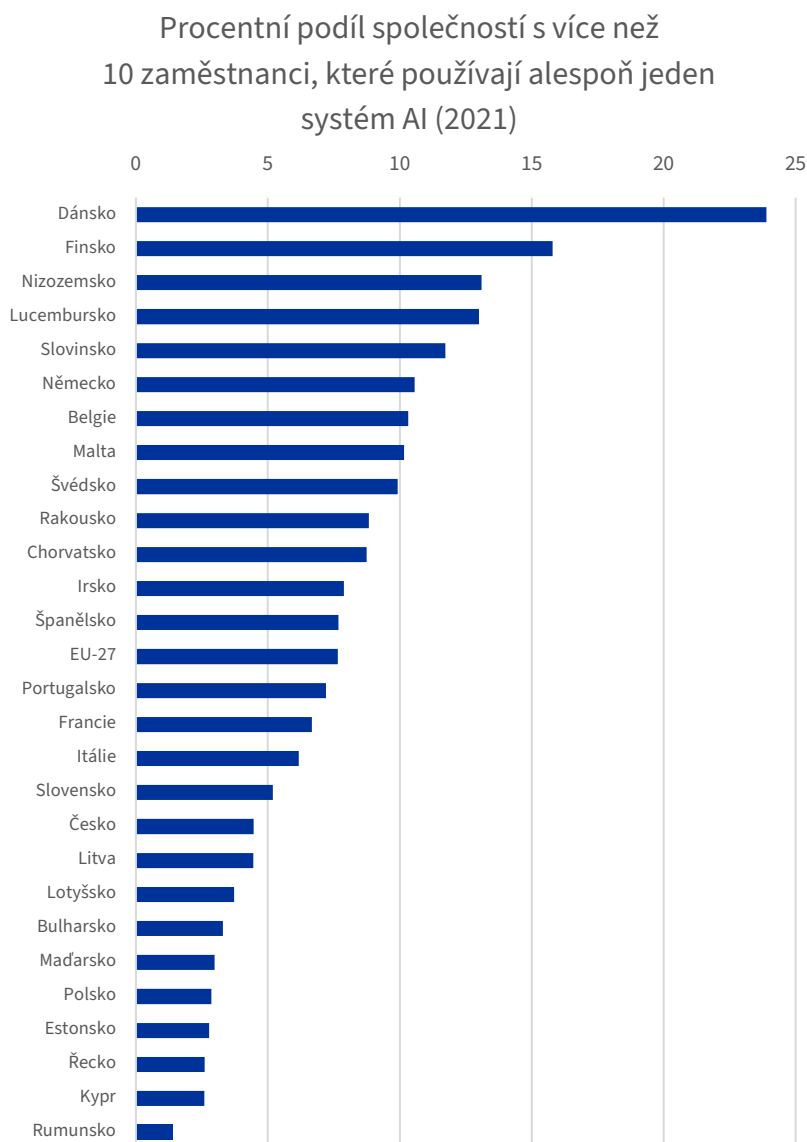
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

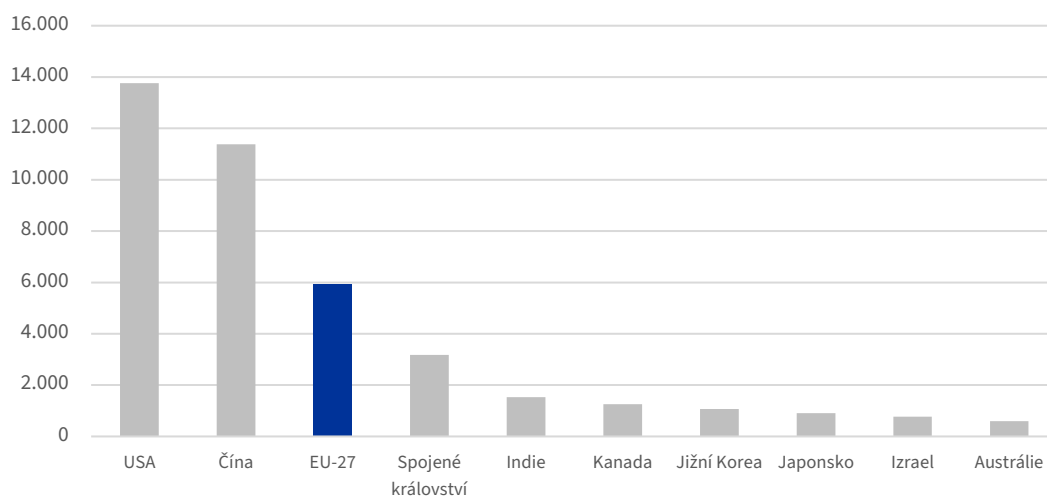
	Evropská unie
Počet obyvatel	447 695 350
HDP na obyvatele (v EUR)	38 150
Míra nezaměstnanosti	6,1 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %



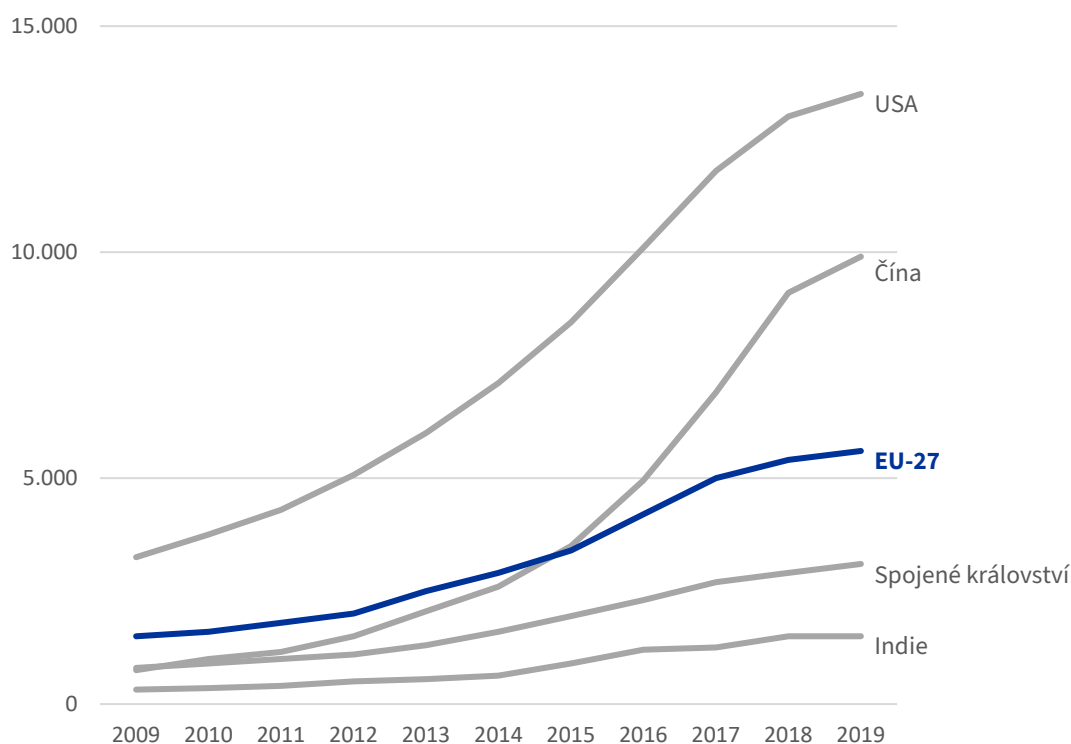
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



RAKOUSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



RAKOUSKO

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Dánsko, Francie, Polsko a Rumunsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vláda vaší země investovala v letech 2012 až 2017 do výzkumu umělé inteligence téměř 350 milionů eur, čímž vytvořila pevný základ pro inovace. Rakouské společnosti nyní stojí v čele vývoje AI, což je přínosné jak pro domácí, tak pro mezinárodní trhy. Podporujete směřování EU v oblasti správy umělé inteligence, ale požadujete větší financování na úrovni EU s cílem dále podpořit výzkum a vývoj.
- Velmi se zasazujete o uplatňování přístupu založeného na účelu prostřednictvím „zásady technologické neutrality“. To znamená, že právní předpisy nesmějí zvýhodňovat ani diskriminovat konkrétní technologie, modely nebo přístupy umělé inteligence. V citlivém odvětví, jako je zdravotní péče, není důležité, zda systém využívá hluboké učení nebo AI založenou na pravidlech – tyto systémy jsou považovány za „vysoce rizikové“, protože v případě selhání nebo chyby na straně umělé inteligence mohou způsobit i smrt člověka.
- Přístup založený na účelu má také zásadní význam pro dovoz systémů umělé inteligence. V mnohých případech se nedají ověřit skutečné schopnosti dovážených systémů pouze na základě prohlášení vývojáře – zejména v případech, kdy zahraniční subjekty mohou mít motivaci skrývat škodlivé prvky, jako je špiónážní software.
- Pokud jde o požadavky na testování, souhlasíte s rámcem preventivního testování, protože podporuje odpovědný přístup a umožňuje snazší sledování příčin možných problémů.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Rakouský technologický institut sehrává ve spolupráci s Interpolem, Europolem a Úřadem OSN pro drogy a kriminalitu se sídlem ve Vídni vedoucí úlohu při integraci systémů AI do národní bezpečnosti prostřednictvím pokročilých sledovacích řešení a platform pro imerzivní odbornou přípravu. Výzkumní pracovníci jsou odhodláni vyvíjet vysoce důvěryhodné a spolehlivé modely umělé inteligence, které přispívají ke zvýšení bezpečnosti Rakouska. Vzhledem k tomu, že národní bezpečnost obvykle nespadá do oblasti působnosti regulačních rámců EU, považujete za rozumné tuto oblast do nařízení o AI nezahrnovat.
- Jak jste již zdůraznili v článku 1, dále je důležité vyjasnit, že oblast působnosti regulačního rámce musí zahrnovat také systémy AI vyvinuté mimo EU a následně dovezené. Nařízení by se mělo vztahovat na všechny systémy AI používané v Evropské unii bez ohledu na jejich původ.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

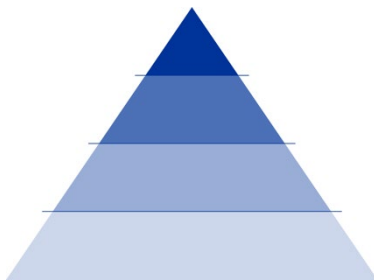
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko	na základě účelu	předběžná opatrnost		národní bezpečnost	
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

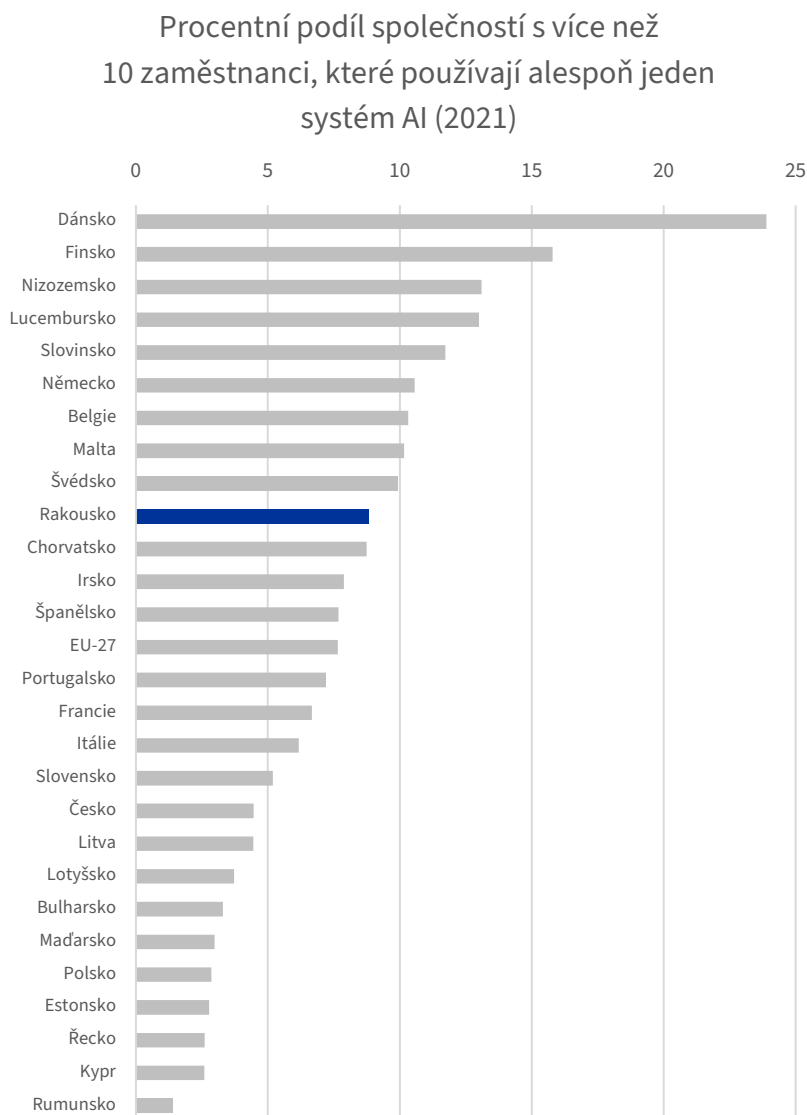
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

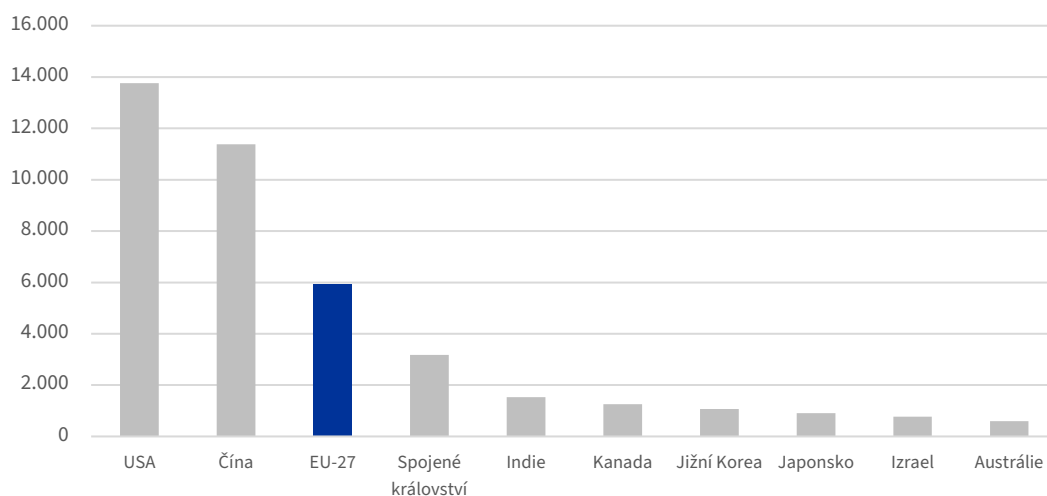
	Evropská unie	Rakousko
Počet obyvatel	447 695 350	9 104 772
HDP na obyvatele (v EUR)	38 150	51 830
Míra nezaměstnanosti	6,1 %	5,1 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	10,8 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	32 %



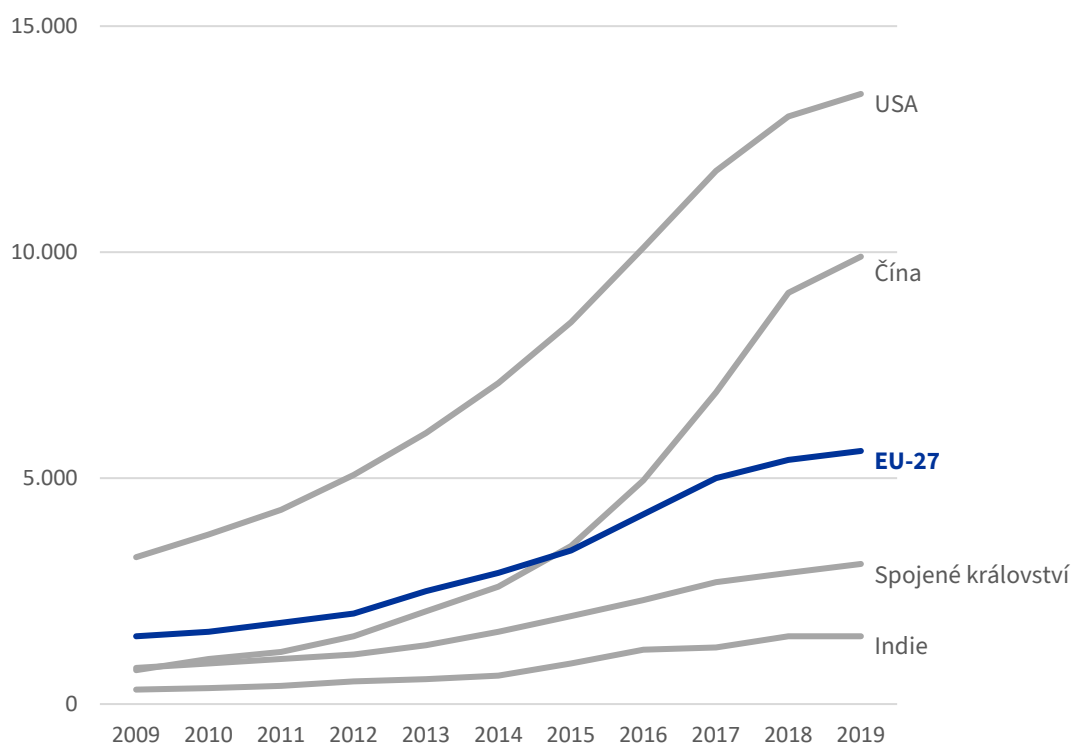
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



BELGIE



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**



BELGIE

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Kypr, Německo, Řecko a Španělsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost do **tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Jednou z hlavních priorit vaší země je prosazování základních práv. Pro řešení vnitrostátních i nadnárodních hrozeb – včetně těch, které vyplývají právě z digitálních nástrojů – je zapotřebí naléhavá přeshraniční spolupráce.
- Chcete zajistit, aby se umělá inteligence vyvíjela odpovědně. Od roku 2010 dosáhli vývojáři významných pokroků a vytvářejí nové nástroje rychleji, než se jim může společnost přizpůsobit. Jen málo jednotlivců dokáže úplně pochopit budoucnost AI: umělá inteligence pro všeobecné účely (GPAI) bude moci brzy napodobovat lidské uvažování, zatímco kvantová umělá inteligence by mohla exponenciálně posílit schopnosti prostřednictvím kvantové výpočetní techniky, což by vedlo k dosud neznámým důsledkům.
- Pokud by se však AI vyvíjela odpovědně, mohla by zlepšit životy a účinněji řešit globální výzvy. Je však také zapotřebí zohlednit etické aspekty. Proto prosazujete co nejbezpečnější a nejodpovědnější přístup. Testování v reálných podmínkách, jak je navrhováno v rámci flexibilního přístupu, není v žádném případě možné. Pro uživatele představuje zbytečná rizika, která lze zmírnit důsledným počátečním testováním na pískovišti.
- Podporujete klasifikaci rizik založenou na schopnostech: jakýkoli systém AI, který se dá zneužít, který může být předpojatý nebo mít škodlivý dopad, by měl být považován za vysoce rizikový. Vývoj AI by měl jít ruku v ruce s podporou rozmanitosti a prosazováním lidských práv, nikoli proti nim.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Zasazujete se o širokou oblast působnosti právních předpisů. Podle vašeho názoru nemá smysl dohodnout se na celounijním nařízení o umělé inteligenci, které by normalizovalo odpovědný vývoj AI, a poté z něj okamžitě odstranit tak důležité oblasti, jako je bezpečnost.
- Jedinou výjimkou, s níž můžete souhlasit, jsou modely s otevřeným zdrojovým kódem. Tyto modely významně urychlují inovace, protože jejich vývoj je jasnější. Vývojáři společně experimentují a vzájemně si vyměňují poznatky, díky čemuž je možné zdokonalovat nástroje. U AI s otevřeným zdrojovým kódem je často snazší provádět přezkumy a kontroly než u modelů s uzavřeným zdrojovým kódem, což je v souladu s cíli aktu.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

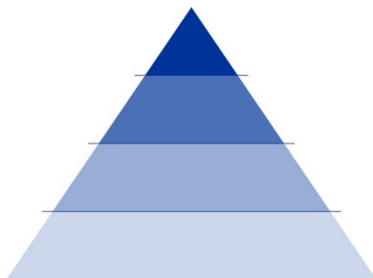
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie	na základě schopností	předběžná opatrnost		otevřený zdroj	
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

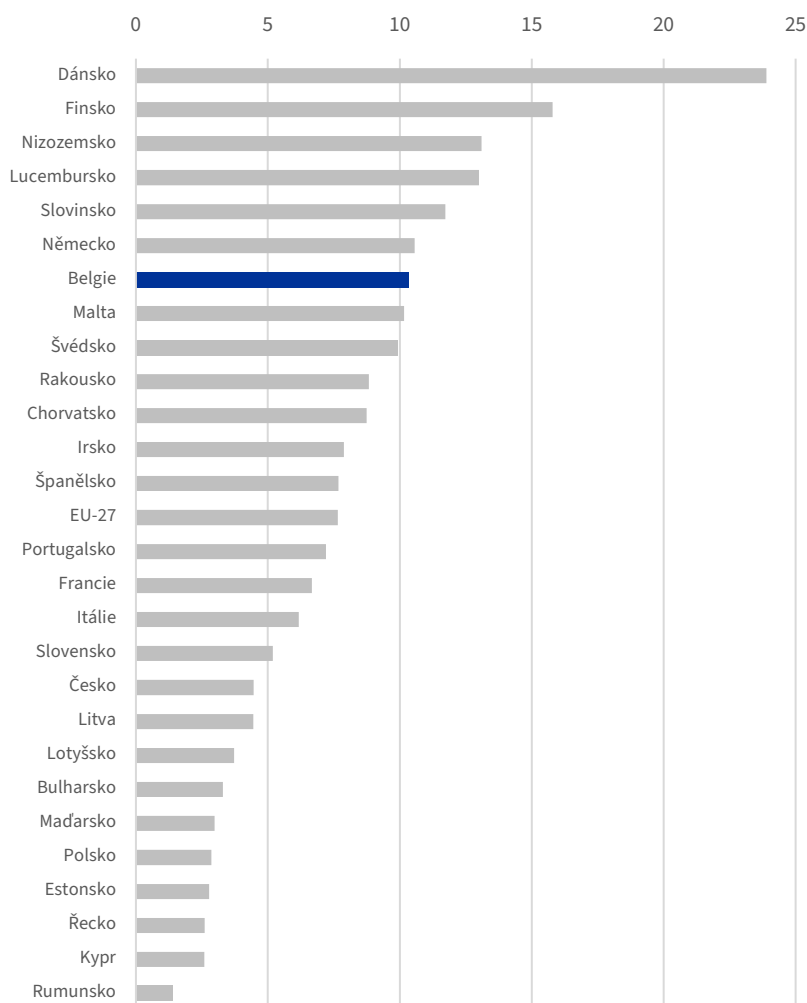
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Belgie
Počet obyvatel	447 695 350	11 742 796
HDP na obyvatele (v EUR)	38 150	50 610
Míra nezaměstnanosti	6,1 %	5,5 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	13,8 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	28,3 %

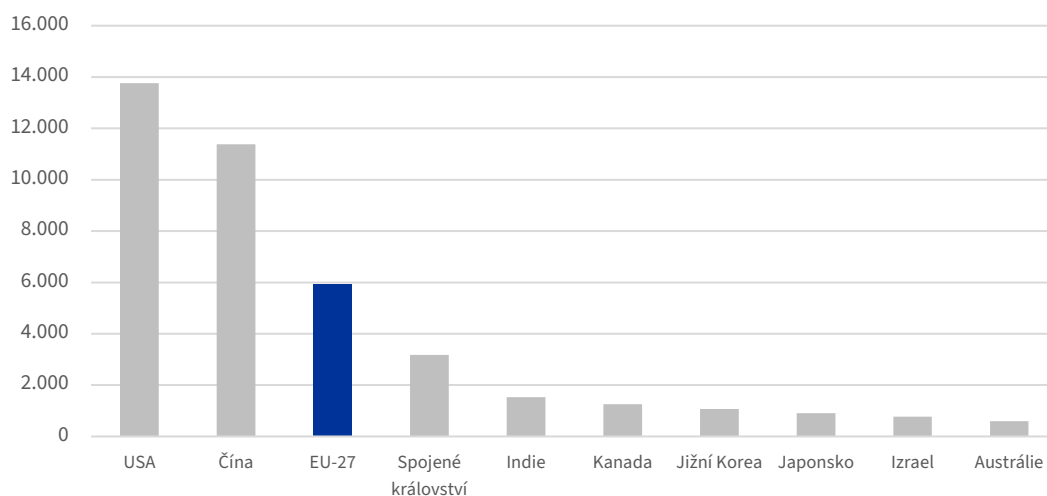
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



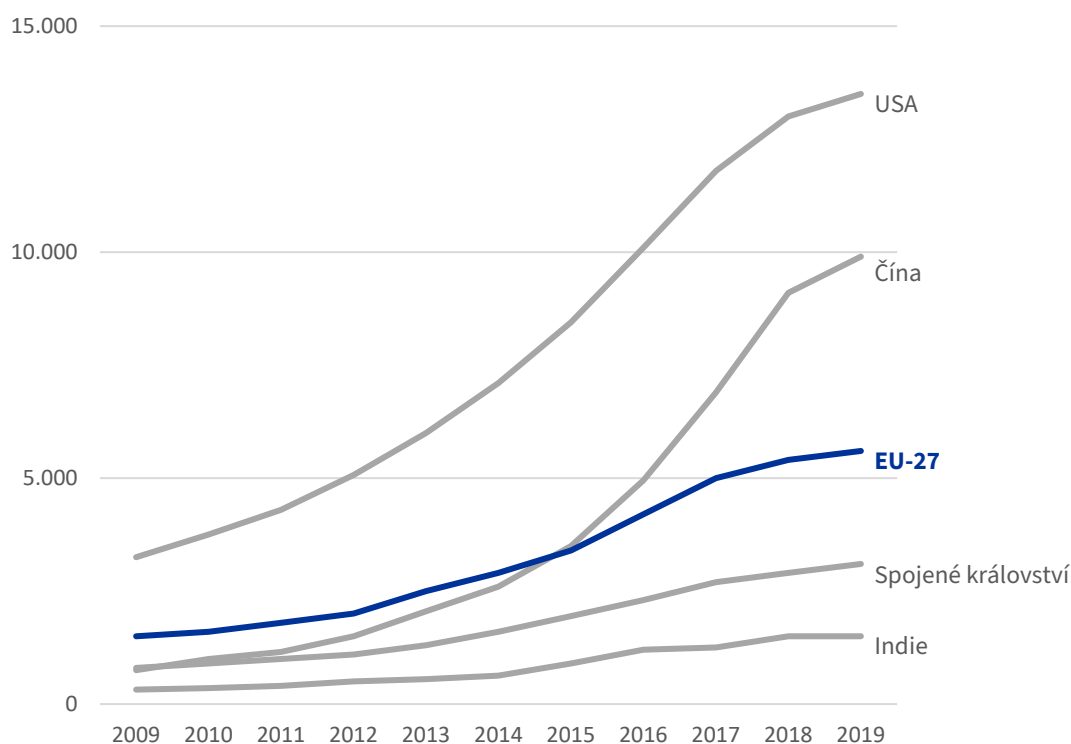
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI

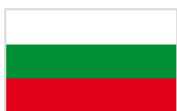


Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



BULHARSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

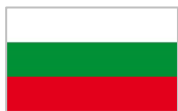
Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.



Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Lucembursko, Malta a Slovensko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost do **tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Bulharskému odvětví špičkových technologií a IT se v oblasti umělé inteligence daří: má silné vazby na klienty v západní Evropě a USA. Umělá inteligence je na bulharském trhu práce již nyní zastoupena 3 % a vy se tento významný podíl snažíte ještě zvýšit. Od strategie EU v oblasti umělé inteligence očekáváte výsledky, které podpoří průmysl a univerzity ve vaší zemi při maximalizaci přínosů AI.
- V zásadě můžete vyjádřit podporu přístupu založenému na účelu. Takový přístup vede vývojáře a provozovatele k tomu, aby pečlivě přemýšleli o tom, jak lze jejich systém využít, a nikoli pouze o tom, co dokáže. To zvyšuje odpovědnost při využívání systému v reálném světě.
- Za důležitější však považujete, aby EU brala v úvahu zeměpisné rozdíly. Požadujete proto zajištění zeměpisné vyváženosti. Bohatší země mají mnohem více finančních i lidských zdrojů na výzkum umělé inteligence, mimo jiné na zdlouhavé a nákladné testovací postupy před spuštěním. Aby byl zajištěn spravedlivý rozvoj v celé EU, potřebuje Bulharsko finanční prostředky na podporu svých malých a středních podniků a startupů.
- Dnes je vaší hlavní prioritou zajištění pevného závazku k financování výzkumu AI v Bulharsku ze strany EU. Jste ochotni diskutovat o požadavcích na řízení rizik a testování, ale až poté, co budou zohledněny vaše zásadní obavy.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- I když podporujete flexibilní právní úpravu uvedenou v článku 1, jste i nadále přesvědčeni, že by měla platit jasná pravidla pro všechna odvětví – včetně národní bezpečnosti a zpravodajských služeb. Udělením výjimek těmto odvětvím vznikají právní mezery a zvyšuje se riziko toho, že nebude odhalena neetická nebo škodlivá AI.
- Cílem není zamezit využívání AI v oblasti národní bezpečnosti, ale zajistit, aby byla AI navržena zodpovědně tak, aby se předešlo předpojatosti nebo diskriminaci.
- Stávající znění návrhu považujete za příliš vágní a rádi byste upřesnili, že stejné nařízení se vztahuje i na modely AI dovážené od vývojářů ze zemí mimo EU. Důvod je prostý: chcete předejít nejasnostem. Účelem aktu o umělé inteligenci je podpořit rozvoj AI v Unii a zvýšit atraktivitu unijních modelů AI v porovnání s americkými nebo čínskými produkty.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

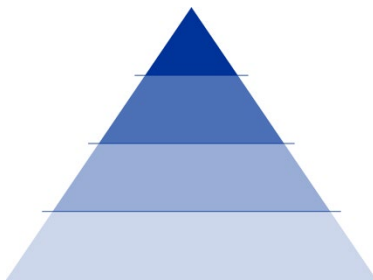
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko	na základě účelu	flexibilní		bez výjimky	
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

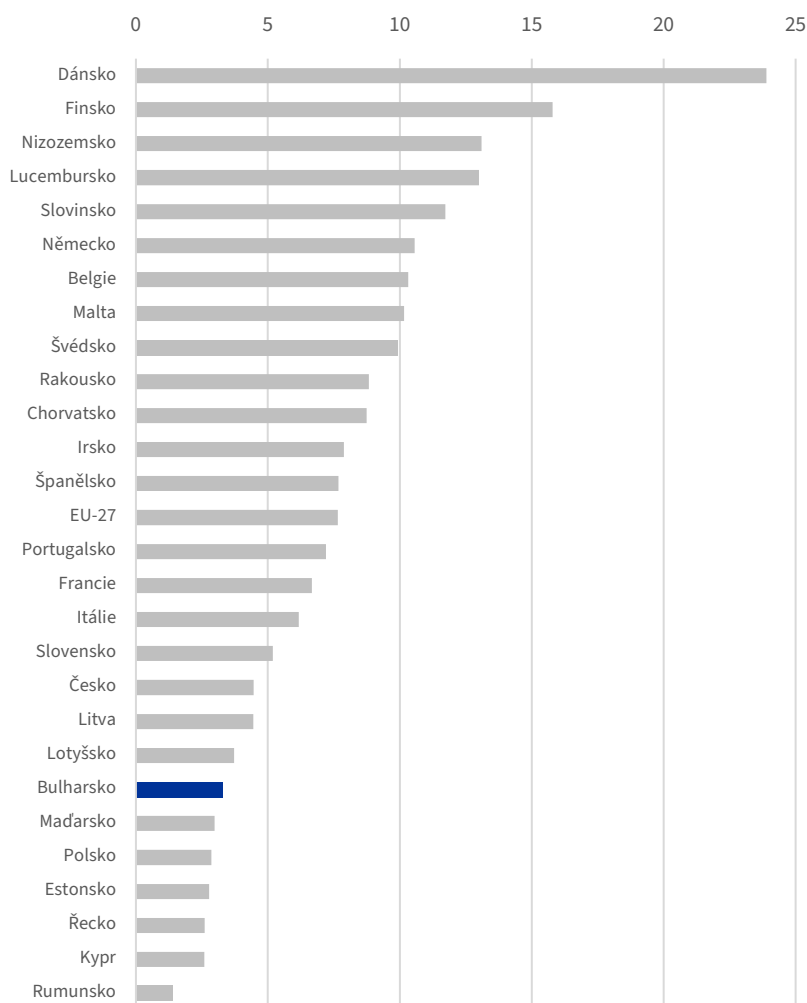
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Bulharsko
Počet obyvatel	447 695 350	6 447 710
HDP na obyvatele (v EUR)	38 150	14 690
Míra nezaměstnanosti	6,1 %	4,3 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	3,6 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	7,7 %

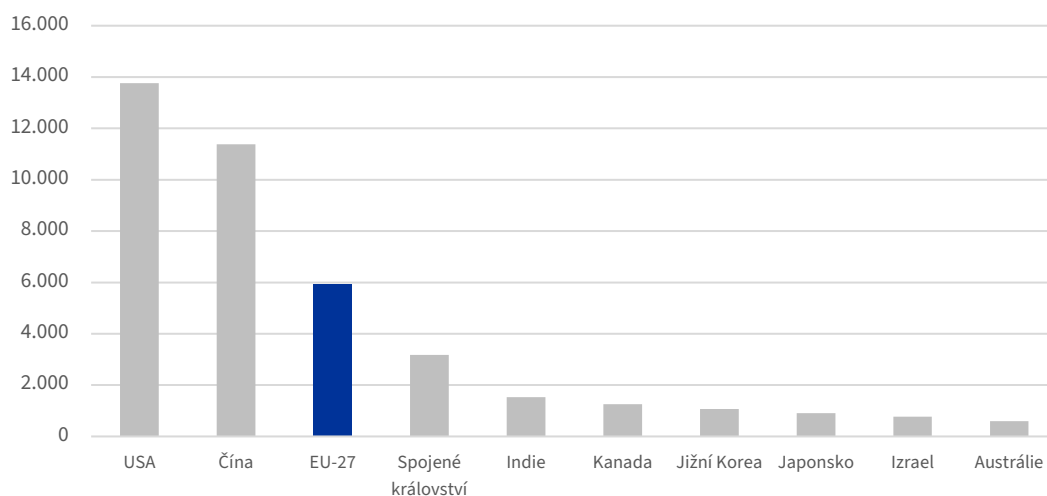
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



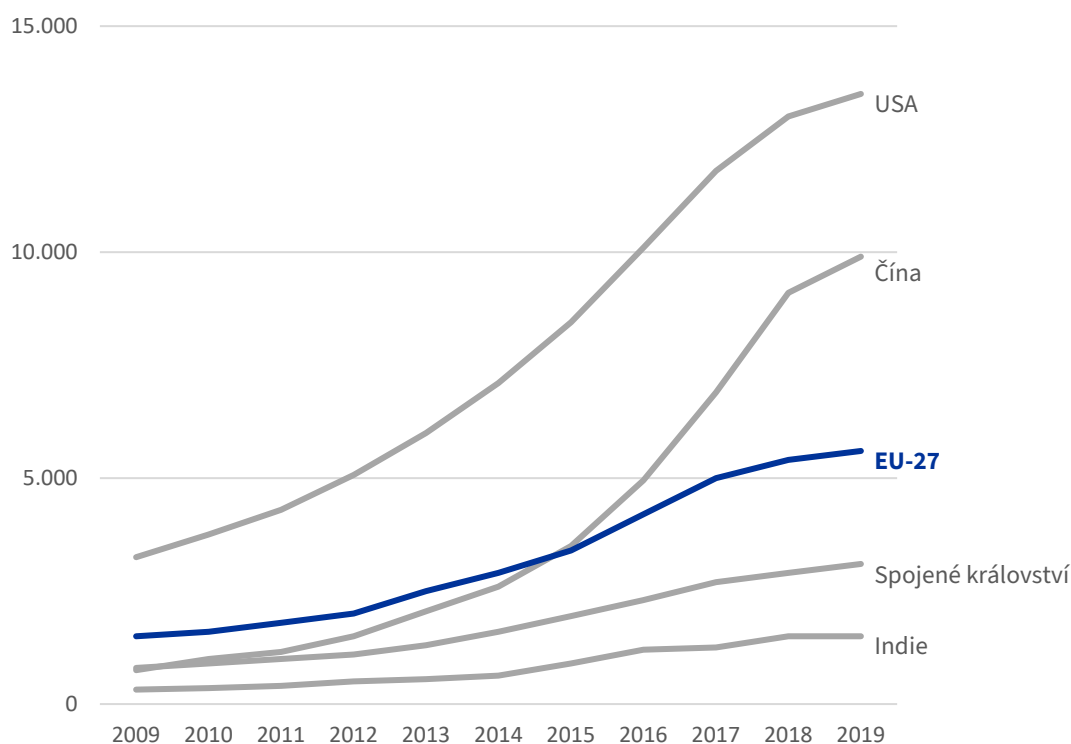
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



KYPR



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



KYPR

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **prečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Belgie, Německo, Řecko a Španělsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost do **tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vláda vaší země upřednostňuje politická opatření, která podporují výzkum a inovace, a zároveň prosazuje přísné etické pokyny. Neregulovaná AI bude nevyhnutelně představovat bezpečnostní hrozbu pro jednotlivce i podniky. Například v roce 2010 způsobily obchodní algoritmy založené na AI propad trhu, kdy během několika minut přišel trh o akcie ve výši 1 bilionu dolarů.
- Podle vašeho názoru musí být při vývoji etické AI upřednostněna transparentnost (aby bylo jasné, jak AI funguje), odpovědnost (aby lidé nadále nesli odpovědnost za rozhodnutí učiněná umělou inteligencí) a mechanismy kontroly (aby se zamezilo fungování umělé inteligence bez kontroly). Máte obavy, že příliš flexibilní pravidla by mohla uvedené zásady oslabit.
- Díky ambiciózním opatřením a přístupu předběžné opatrnosti ve vztahu k testování, který zahrnuje jasná a vymahatelná pravidla, se EU může stát odpovědným a progresivním lídrem na poli inovací v oblasti AI.
- Podporujete klasifikaci rizik založenou na schopnostech: jakýkoli systém AI, který se dá zneužít, který může být předpojatý nebo mít škodlivý dopad, by měl být považován za vysoce rizikový – ne proto, že by umělá inteligence byla ze své podstaty nebezpečná, ale kvůli možnostem jejího využití. Jste však **ochotni diskutovat o klasifikaci rizik založené na účelu, pokud bude za dva roky klasifikace přezkoumána.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Jste pevně přesvědčeni, že vývoj etické umělé inteligence se musí uplatňovat plošně – žádné odvětví ani žádný vývojář by neměli být vyňati z oblasti působnosti nařízení EU o umělé inteligenci. Právě v oblastech, jako je vymáhání práva, armáda, obrana a národní bezpečnost, je riziko zneužití, předpojatosti a porušování práv nejvyšší. Pro tato odvětví by neměly platit žádné výjimky, měly by se na ně naopak vztahovat nejpřísnější normy.
- EU chce být světovým lídrem v oblasti důvěryhodné a odpovědné AI. Proto se zasazujete o to, aby šla v tomto ohledu příkladem. Přehlížení určitých problematických aplikací by oslabilo důvěryhodnost a účinnost aktu EU o umělé inteligenci. Při používání AI – zejména v oblastech, které mohou mít největší dopad na životy a svobody lidí – je třeba se řídit etickými zásadami.
- Článek 2 je pro vás důležitější než článek 1. Pokud se vám nepodaří přesvědčit ostatní o svých zájmech týkajících se článku 1, zaměřte se na to, abyste prosadili vaše zásadní požadavky vztahující se k článku 2.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

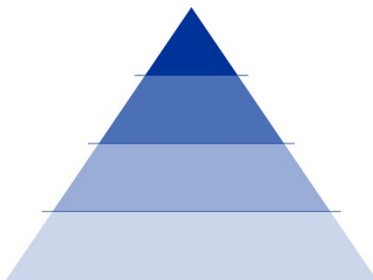
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr	na základě schopností	předběžná opatrnost	přezkum za 2 roky	bez výjimky	
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

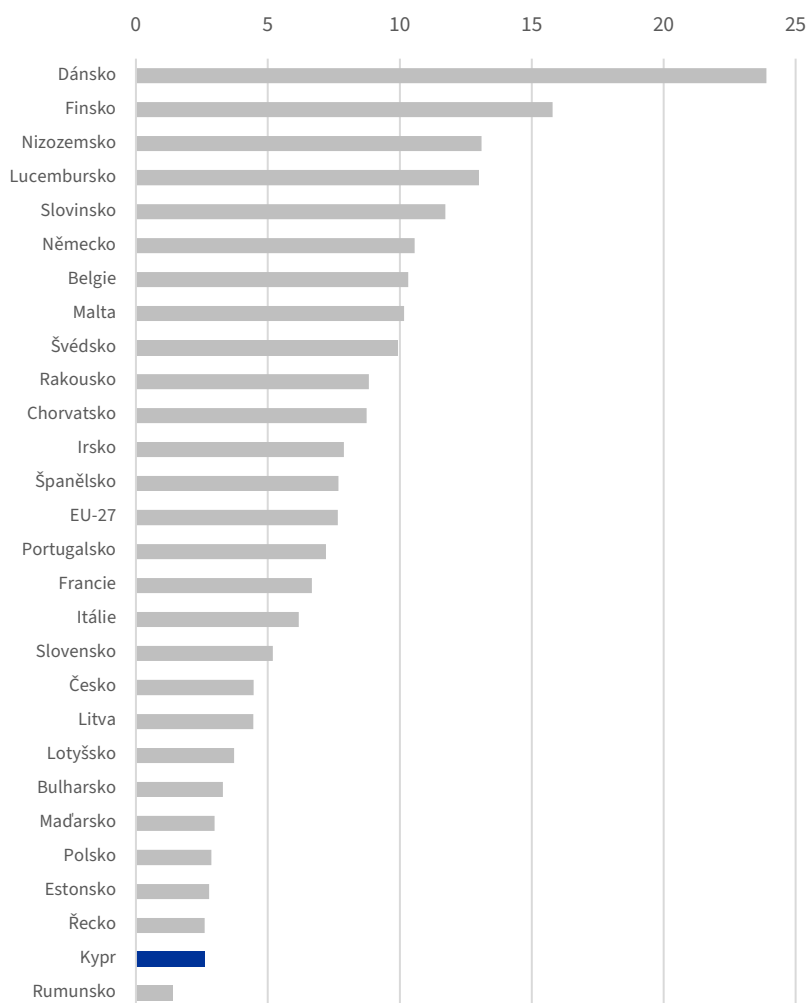
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Kypr
Počet obyvatel	447 695 350	949 084
HDP na obyvatele (v EUR)	38 150	32 720
Míra nezaměstnanosti	6,1 %	5,8 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	4,7 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	25 %

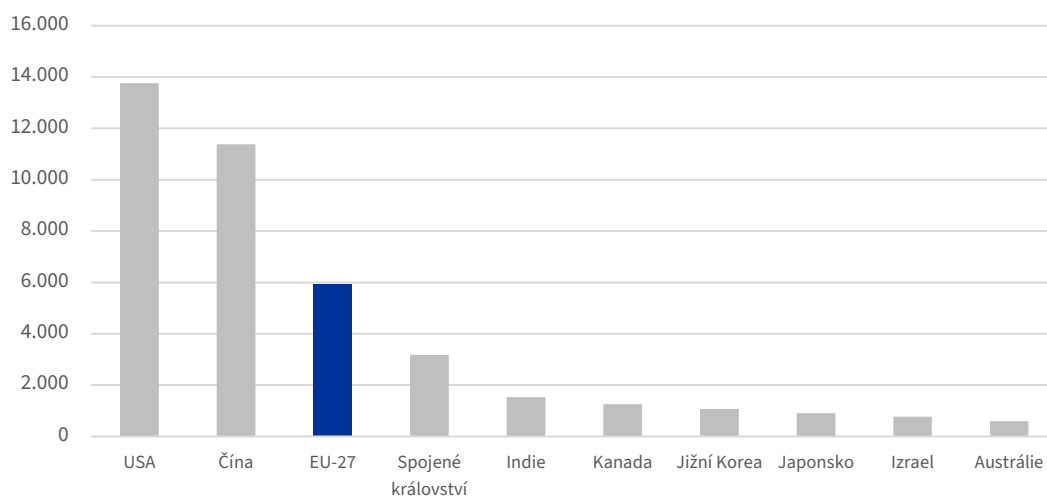
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



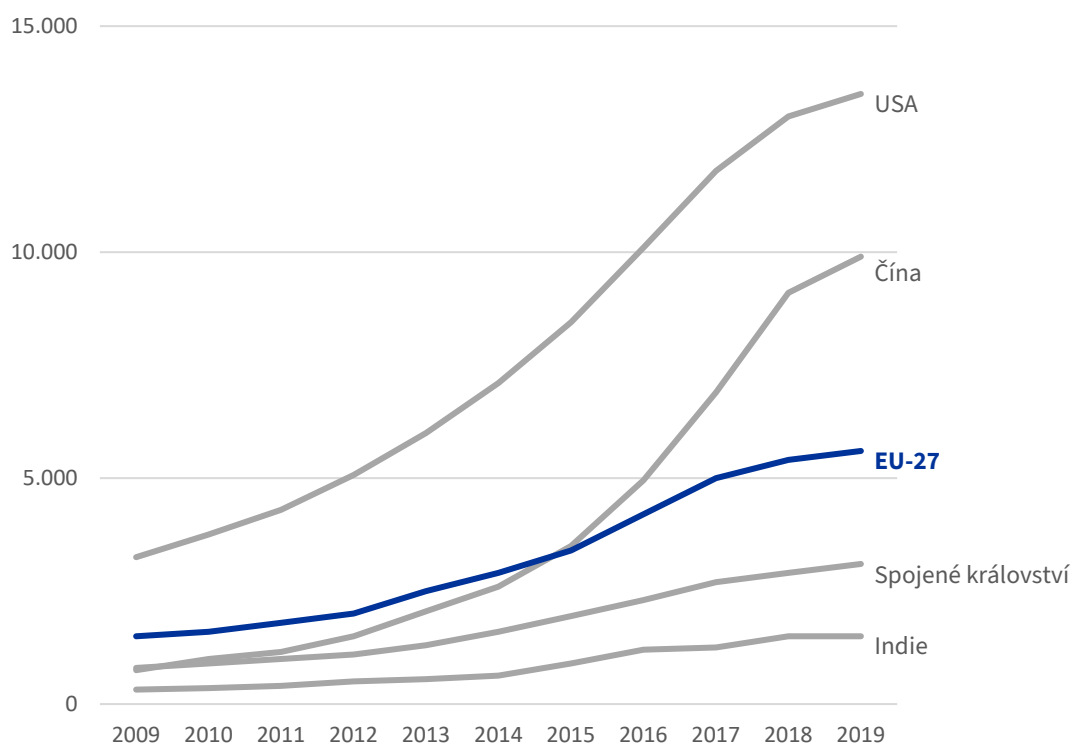
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

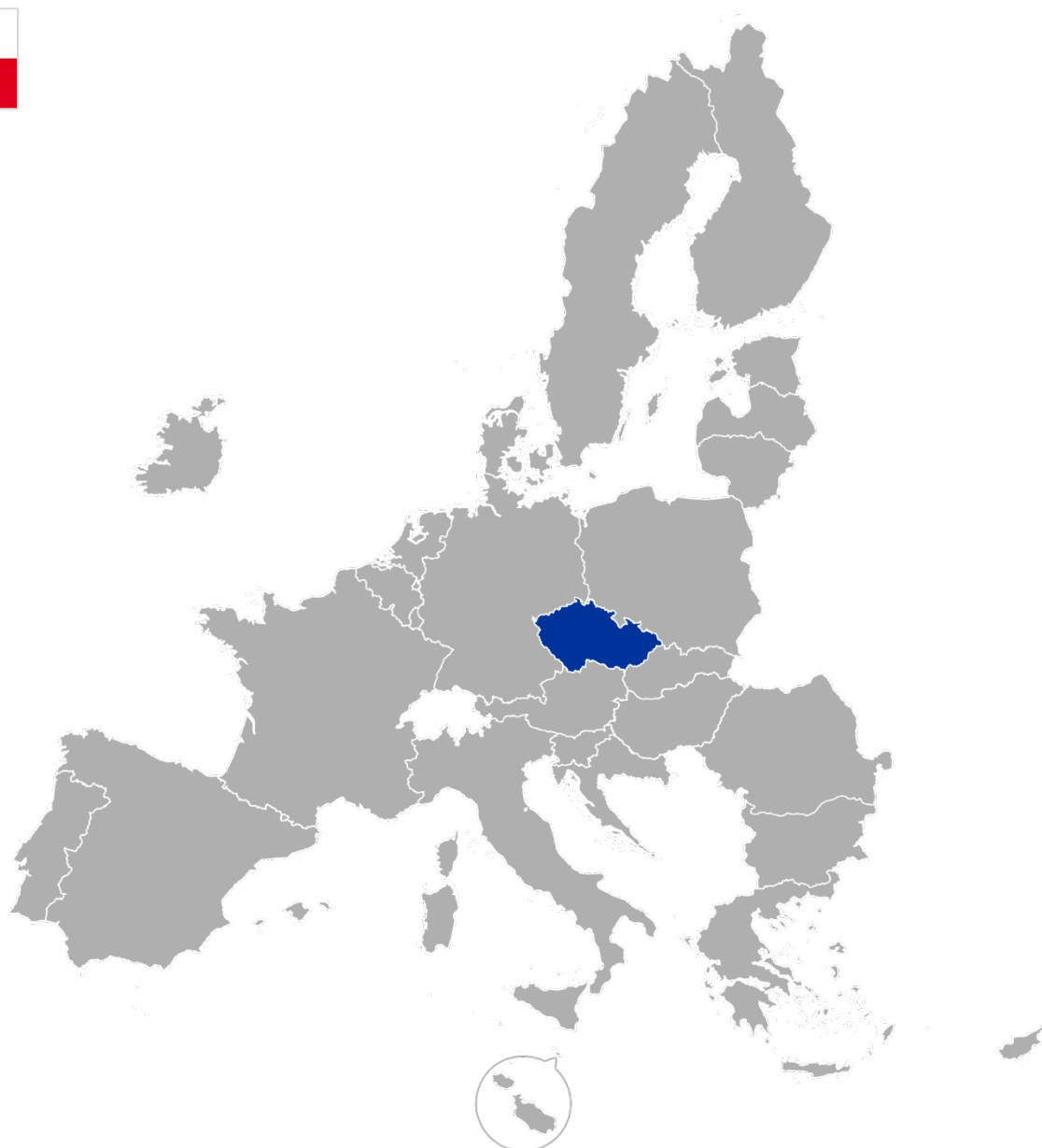
Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



ČESKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



ČESKO

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **prečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Estonsko, Litva, Lotyšsko a Maďarsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je pro přehlednost do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vláda vaší země usiluje o podporu inovací v oblasti AI. Přestože má EU silné postavení v oblasti výzkumu, světový trh řídí USA a Čína, které v současné době určují pravidla. Rychle postupují i další země, jako jsou Indie, Kanada a Spojené království.
- Podle vás z toho logicky vyplývá, že EU musí vytvořit vyvážený regulační rámec, který nebude příliš striktní, aby podnítila inovace v oblasti AI a stala se v této oblasti světovým lídrem. Akt EU o umělé inteligenci by tak mohl přinést zásadní změnu, pokud by získal důvěru a podporu investorů i vývojářů.
- Podporujete přístup založený na schopnostech, neboť přístup založený na účelu může být v rozporu se stávajícími vnitrostátními a ujednáními odvětvovými politikami a vyvolat nejistotu. Klasifikovat všechny modely umělé inteligence pro všeobecné účely (GPAI) jako „vysoce rizikové“ by však bylo příliš přísné. Mnoho systémů GPAI, jako je například umělá inteligence používaná v oblasti zákaznických služeb nebo překladů v reálném čase, představuje nízké riziko. Nadměrná regulace by mohla odrazovat od inovací a investic a posilovat představu, že „EU má nejpřísnější pravidla pro AI“, což by v konečném důsledku oslabilo konkurenceschopnost Evropy.
- Pokud se dnes ukáže, že dosáhnout kompromisu je obtížné, **jste ochotni přidat do nařízení ustanovení o přezkumu za pět let**, čímž se zajistí, že v rychle se vyvíjejícím prostředí AI zůstane nařízení stále relevantní.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Jste pevně přesvědčeni, že vývoj etické umělé inteligence se musí uplatňovat plošně – žádné odvětví ani žádný vývojář by neměli být vyňati z oblasti působnosti nařízení EU o umělé inteligenci. Právě v oblastech, jako je vymáhání práva, armáda, obrana a národní bezpečnost, je riziko zneužití, předpojatosti a porušování práv nejvyšší. Pro tato odvětví by neměly platit žádné výjimky, měly by se na ně naopak vztahovat nejpřísnější normy.
- Pokud se chce EU stát světovým lídrem v oblasti důvěryhodné a odpovědné AI, musí jít jasným příkladem. Přehlížení určitých problematických aplikací by oslabilo důvěryhodnost a účinnost aktu EU o umělé inteligenci. Při používání AI – zejména v oblastech, které mohou mít největší dopad na životy a svobody lidí – je třeba se řídit etickými zásadami.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

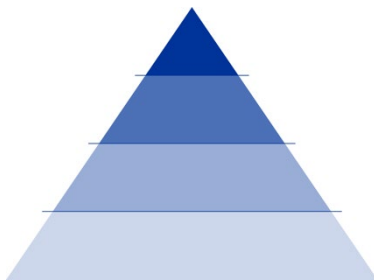
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko	na základě schopností	flexibilní	přezkum za 5 let	bez výjimky	
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

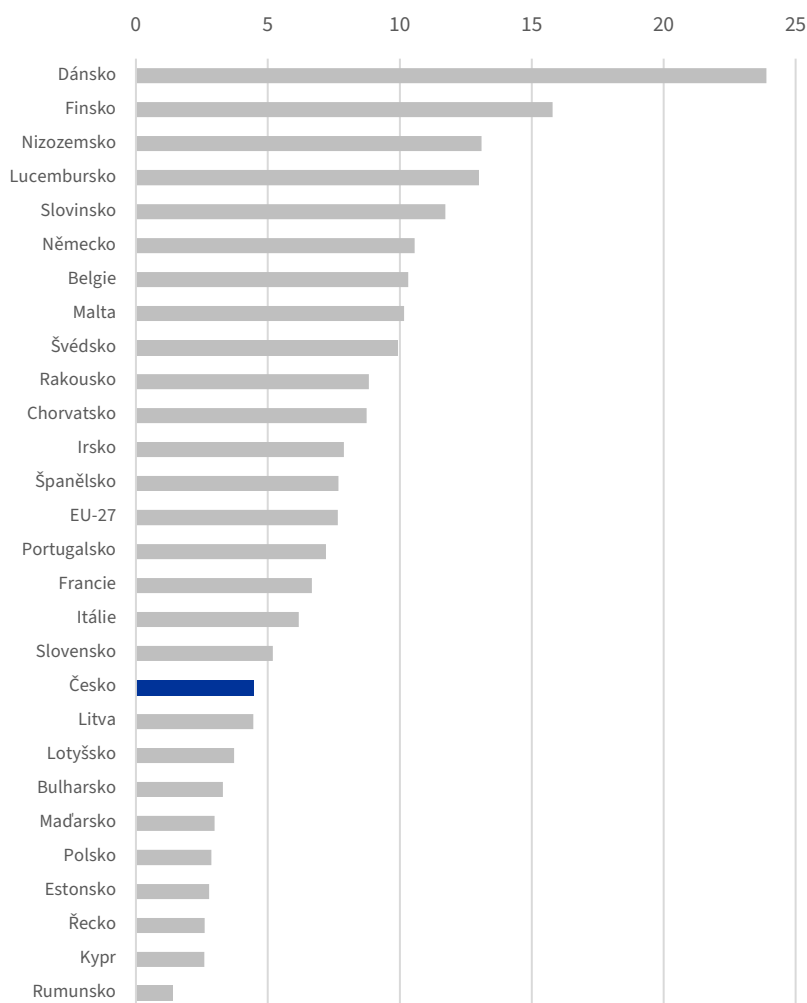
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Česko
Počet obyvatel	447 695 350	10 827 529
HDP na obyvatele (v EUR)	38 150	29 180
Míra nezaměstnanosti	6,1 %	2,6 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	5,9 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	35,5 %

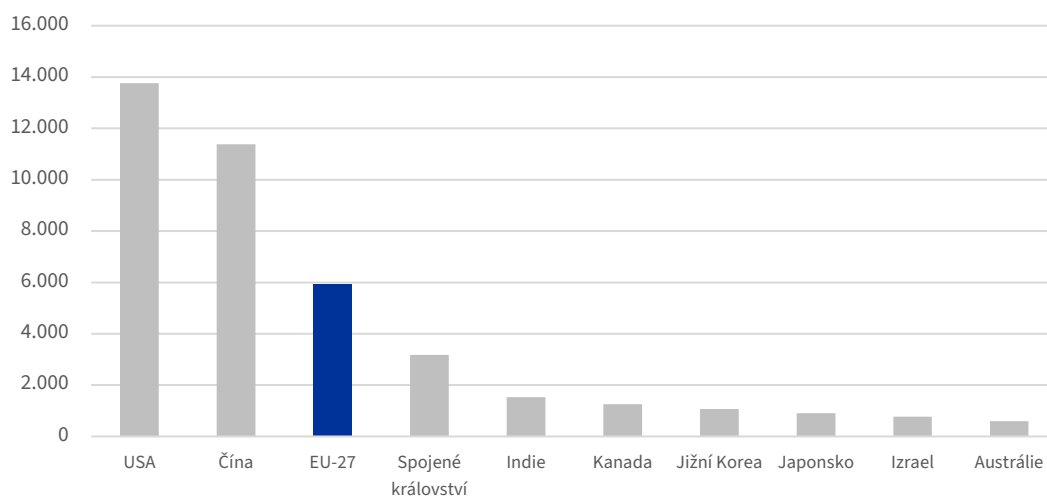
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



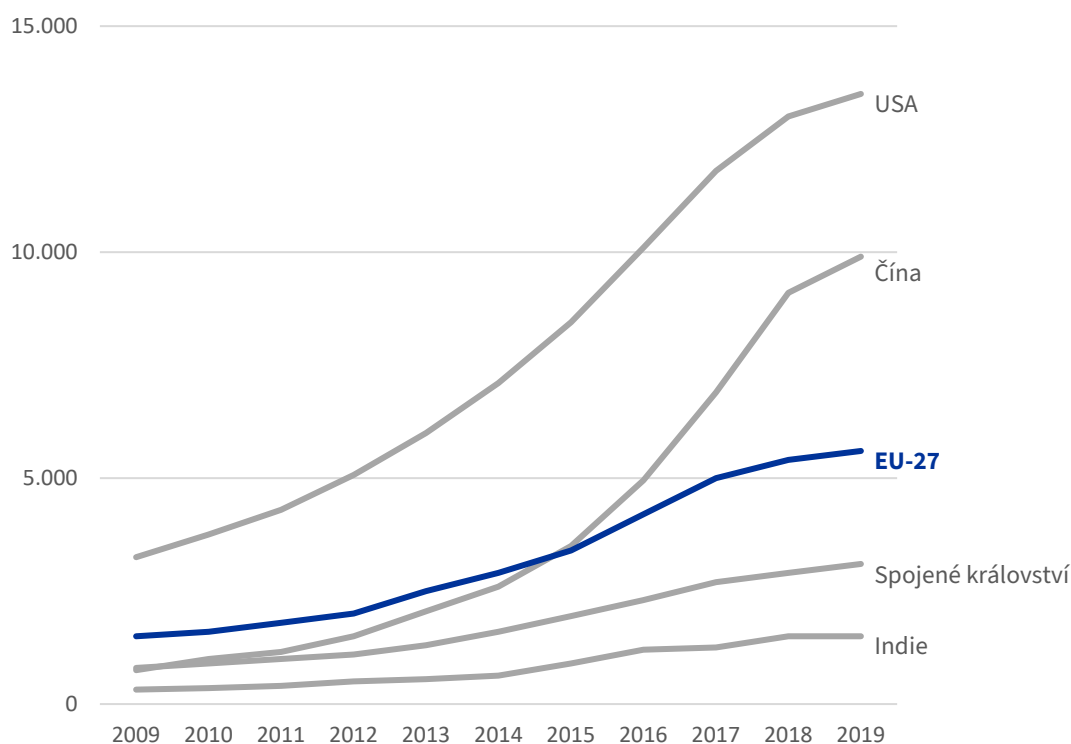
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

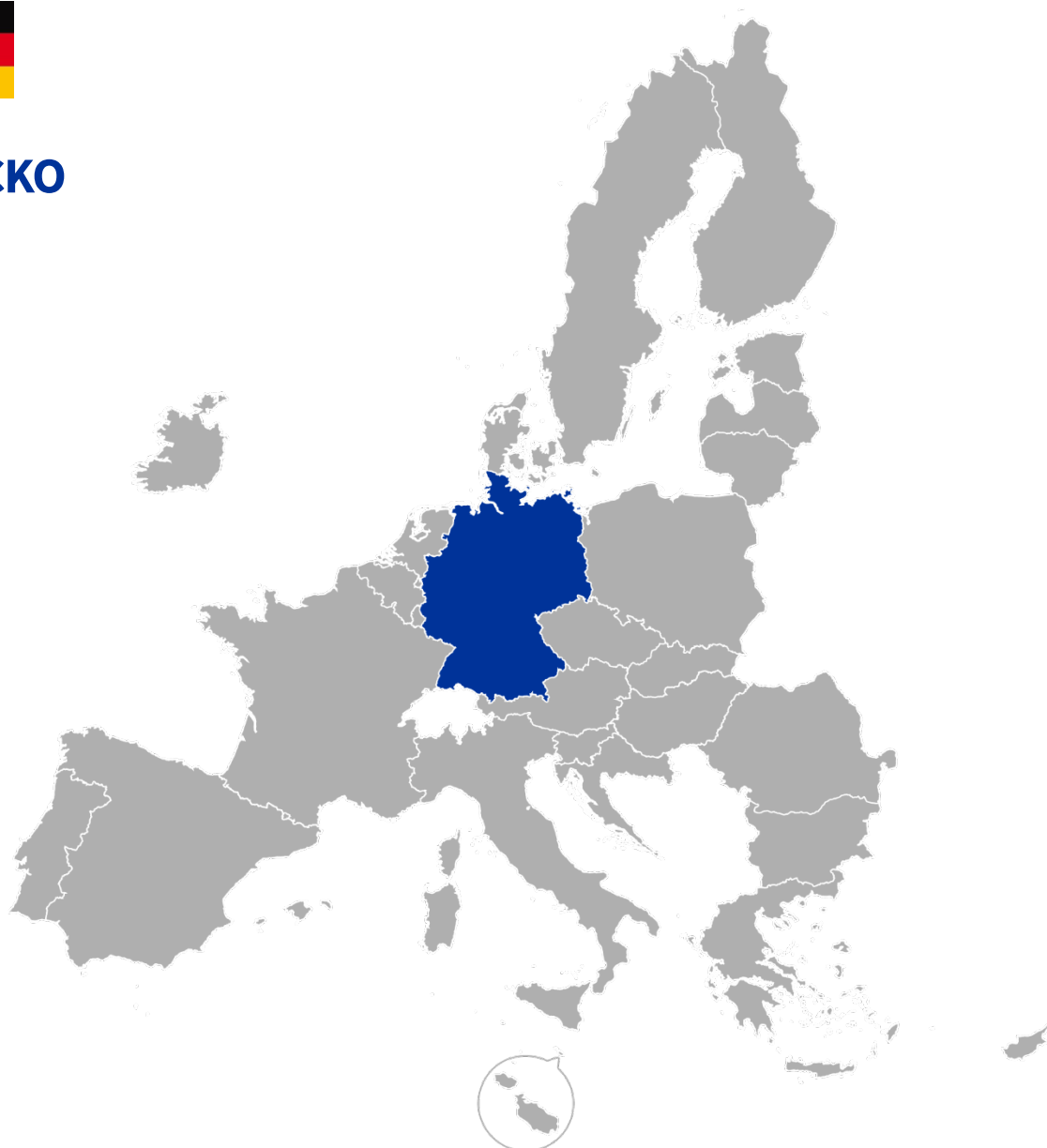
Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



NĚMECKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



NĚMECKO

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílají země, jako je Belgie, Kypr, Řecko a Španělsko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost do **tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vaší hlavní prioritou je prosazování základních práv a uznáváte, že k řešení vnitrostátních i nadnárodních hrozeb, včetně hrozeb vyplývajících z digitálních technologií, je naléhavě zapotřebí přeshraniční spolupráce. Byli jste svědky toho, jak může AI přispívat k dezinformačním a deepfake kampaním, což představuje vážné riziko pro svobodu projevu i demokratické instituce.
- I když se AI vyvíjí rychle a bude hrát ve společnosti stále významnější úlohu, musí být podle vás koncipována tak, aby byla pro lidi a životní prostředí přínosná a nikoli, aby jim škodila.
- Jste pro zavedení přísných požadavků na testování vysoce rizikových modelů, protože hlavním problémem systémů AI je zaujatost. Pokud se AI školí na nereprezentativních souborech dat, která nejsou dostatečně různorodá, může generovat diskriminační výsledky. Tento případ již nastal v případě nástroje společnosti Amazon pro nábor pracovníků, který byl zrušen, protože na základě minulých modelů náborů upřednostňoval mužské kandidáty. Dnes je vaším hlavním cílem vytvořit rámec, který eliminuje i sebemenší riziko újmy způsobené umělou inteligencí, jako jsou zavádějící informace. Víte, že je to ambiciózní cíl, ale považujete jej za nezbytný.
- Podporujete klasifikaci rizik založenou na schopnostech: jakýkoli systém AI, který se dá zneužít nebo může mít škodlivý dopad, by měl být považován za vysoce rizikový – ne proto, že by umělá inteligence byla ze své podstaty nebezpečná, ale kvůli možnostem jejího využití.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Velmi vás znepokojuje, že některé země chtějí, aby se akt EU o umělé inteligenci nevztahoval na vymáhání práva, armádu, obranu nebo národní bezpečnost. Právě v těchto oblastech může AI prostřednictvím diskriminace, falešně pozitivních výsledků a předpojatosti způsobit vážnou újmu a právě v těchto oblastech jsou nejvíce zapotřebí silné záruky. Například zkreslené rozpoznávání obličejů při policejní práci již vedlo k neoprávněnému zatýkání nepřiměřeného počtu mladých mužů na základě rasových předsudků a k závažnému porušování práv jednotlivců.
- Chápete, že v mnoha zemích dochází díky využívání umělé inteligence v těchto odvětvích k velkému zvýšení efektivity, a tento přínos uznáváte. Lidská práva však nelze brát na lehkou váhu. Pozorně sledujete vývoj AI v USA a Číně a velmi vás znepokojuje nedostatek silné ochrany v jejich systémech.
- Smyslem aktu EU o umělé inteligenci je podpora odpovědné AI zaměřené na člověka. Stanovení výjimek pro určitá odvětví by tento cíl oslabilo a podkopalo by ambice EU stanovit globální standard pro důvěryhodnou AI.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

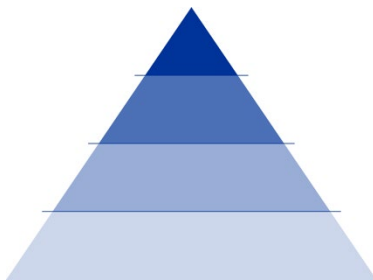
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabíjet bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo	na základě schopností	předběžná opatrnost		bez výjimky	
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

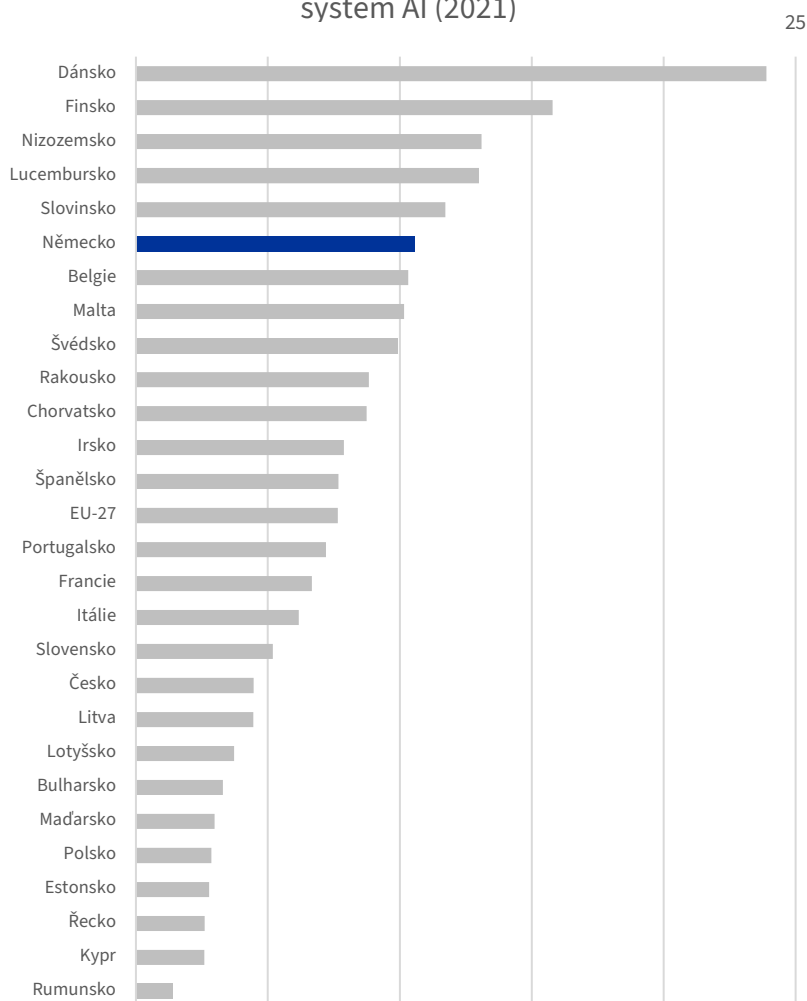
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Německo
Počet obyvatel	447 695 350	83 118 501
HDP na obyvatele (v EUR)	38 150	49 520
Míra nezaměstnanosti	6,1 %	3,1 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	11,6 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	19,8 %

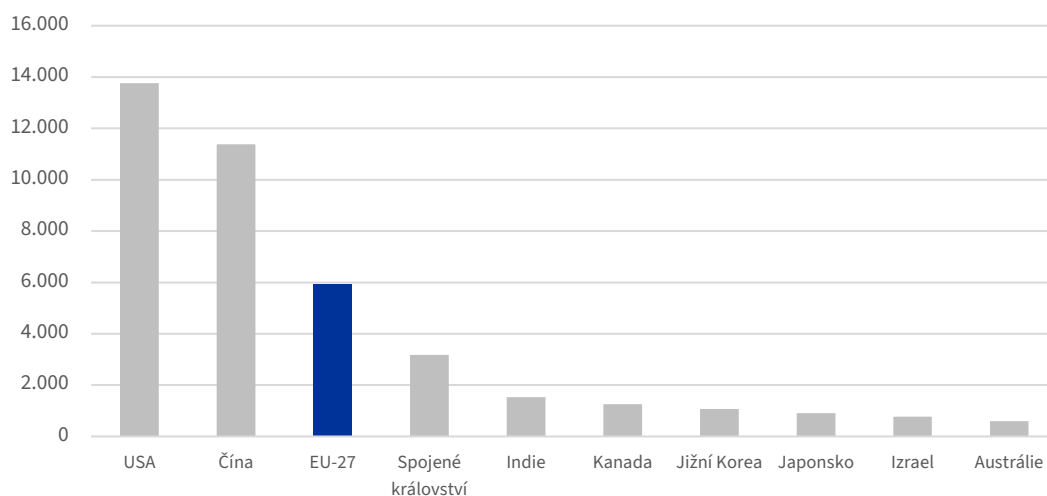
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



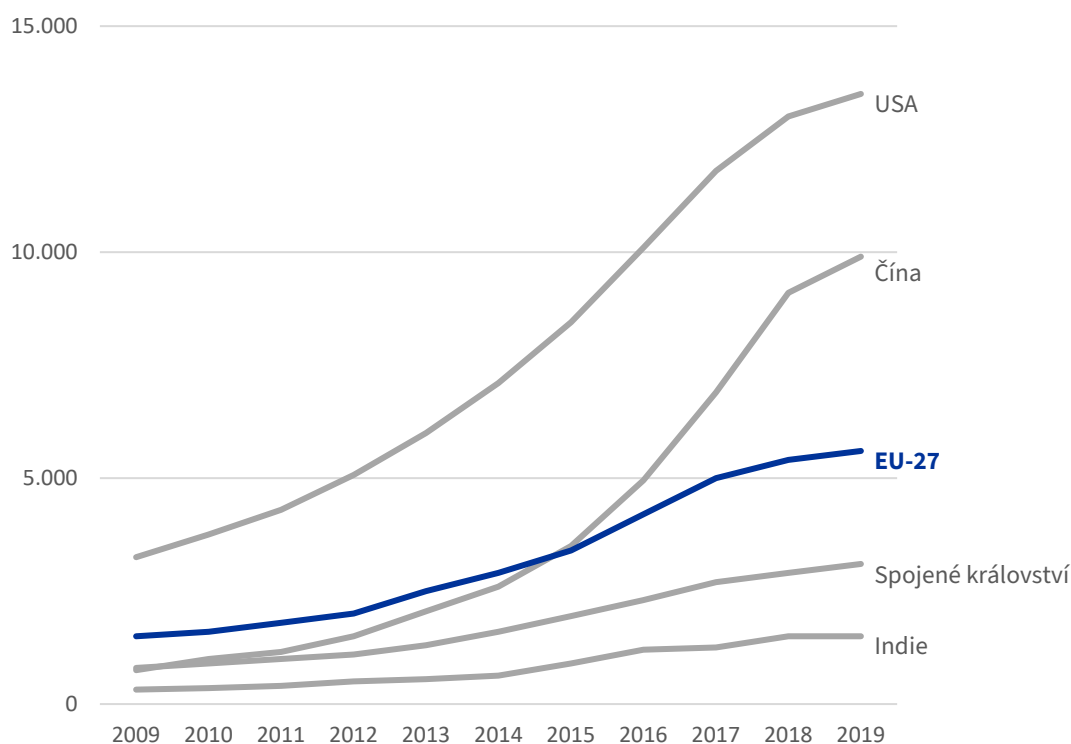
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

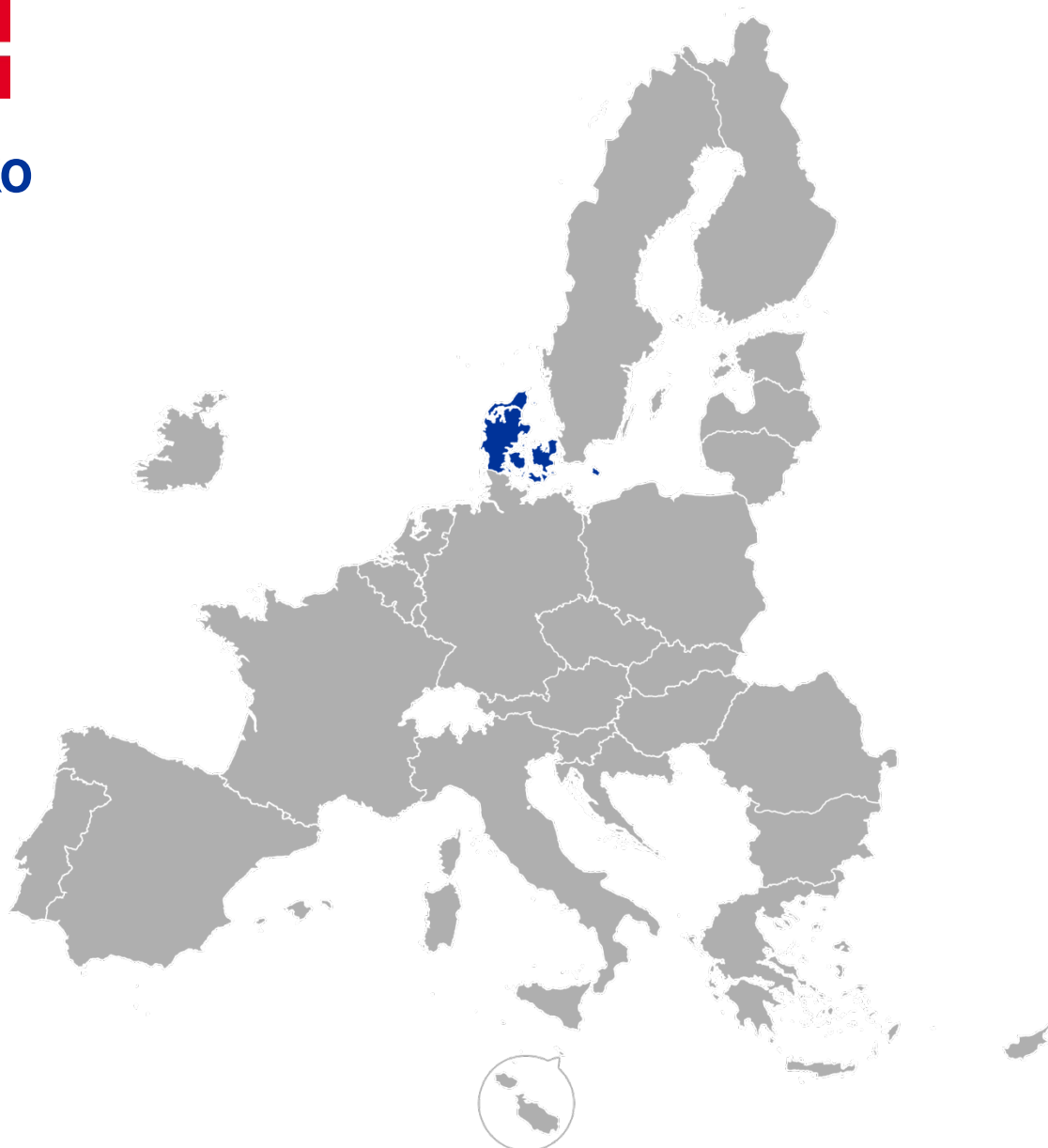
Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



DÁNSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se zájmů a postojů vaší skupiny.**
2. **Některé z vašich argumentů s vámi sdílají země, jako je Francie, Polsko, Rakousko a Rumunsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je pro přehlednost do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vítáte návrh Komise a jako konstruktivní první krok k řešení problémů spojených s vývojem AI vnímáte přístup založený na účelu. Přejete si však do oblasti působnosti zahrnout oblast vzdělávání. Pokud školy používají k monitorování chování, známkování prací nebo předvídání výsledků či výkonu umělou inteligenci, mohou mít studenti pocit, že je stroj neustále pozoruje nebo hodnotí. Mohou tak být vystaveni tlaku, aby vždy „podávali určitý výkon“, i když zrovna mají špatný den. Na rozdíl od učitele z masa a kostí AI často postrádá empatii nebo pochopení pro osobní problémy studentů.
- AI přináší nové způsoby rozhodování, což vyvolává etické otázky. Z tohoto důvodu jste pro zavedení rámce pro preventivní testování, který prosazuje kulturu odpovědnosti a usnadňuje zjišťování příčin problémů, kdykoli se nějaké vyskytnou.
- Kromě toho musí společný rámec EU zohledňovat jazykovou rozmanitost Unie. Většina velkých modelů AI je trénována na datech z široce používaných jazyků, jako je angličtina, francouzština nebo němčina, což může vést k nízké výkonnosti AI u méně rozšířených jazyků, jako je dánština.
- **Pokud se EU rozhodne vyčlenit finanční prostředky na podporu rozvoje AI, bylo by vhodné, aby tato podpora byla prioritně zaměřena na menší země a méně rozšířené jazyky.** Bez toho by vývoj AI mohl prohloubit nerovnosti a vést k poskytování služeb nižší kvality právě v méně rozšířených jazycích.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Nadále vnímáte, že je zásadní, aby do oblasti působnosti nařízení o AI nespadaly záležitosti týkající se národní bezpečnosti, armády nebo obrany. Členské státy jsou odhodlány zachovat si kontrolu nad rozhodnutími v těchto oblastech, které stejně za normálních okolností nespadají do pravomoci EU. To se odráží i ve skutečnosti, že Unie (zatím) nemá jednotnou armádu EU.
- Dále argumentujete tím, že systémy AI používané pro shromažďování zpravodajských informací, kybernetickou obranu nebo plány vojenských operací nemohou podléhat stejným požadavkům na zveřejnění nebo regulačním požadavkům jako civilní systémy, protože by tak mohlo být ohroženo utajení operací, zájmy národní bezpečnosti a strategická výhoda, kterou mají tyto systémy poskytovat.
- Pokud vaše požadavky nelze v dnešní diskusi plně zohlednit, jste ochotni přistoupit na kompromis spočívající v tom, že AI pro vojenské účely a účely národní bezpečnosti nebude z působnosti nařízení vyloučena, ale bude považována za AI s nízkým rizikem, na kterou by se vztahovaly mírnější požadavky.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

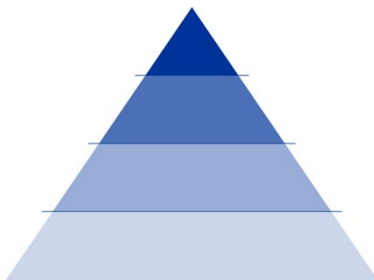
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko	na základě účelu	předběžná opatrnost	finanční podpora	armáda/obrana + národní bezpečnost	
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

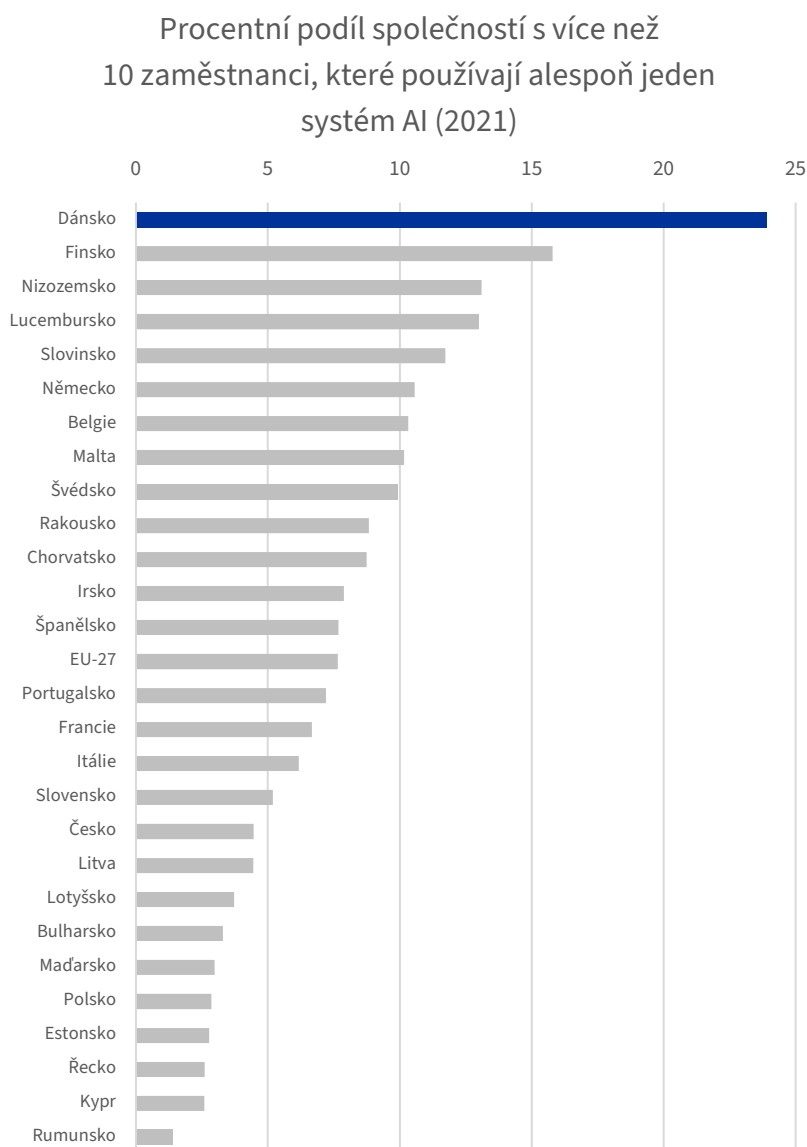
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

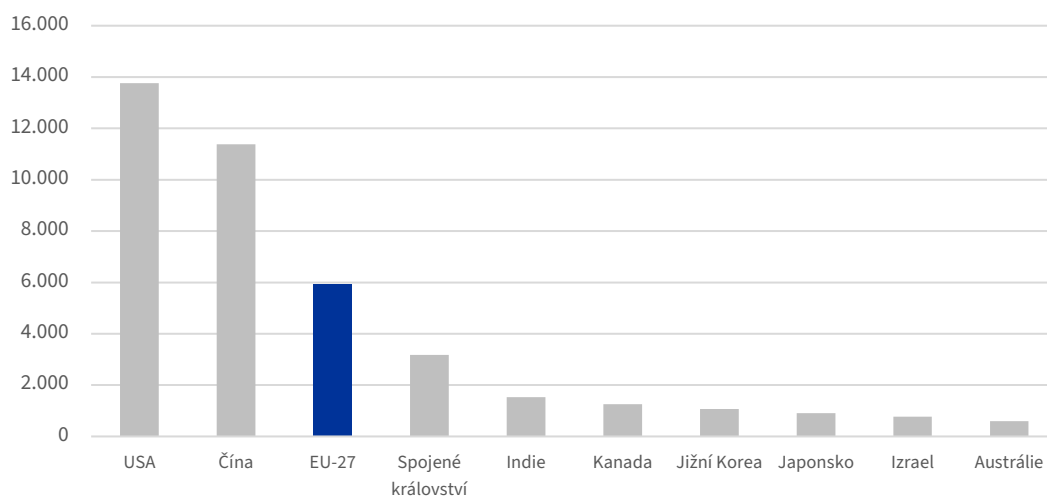
	Evropská unie	Dánsko
Počet obyvatel	447 695 350	5 932 654
HDP na obyvatele (v EUR)	38 150	63 290
Míra nezaměstnanosti	6,1 %	5,1 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	15,2 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	39,4 %



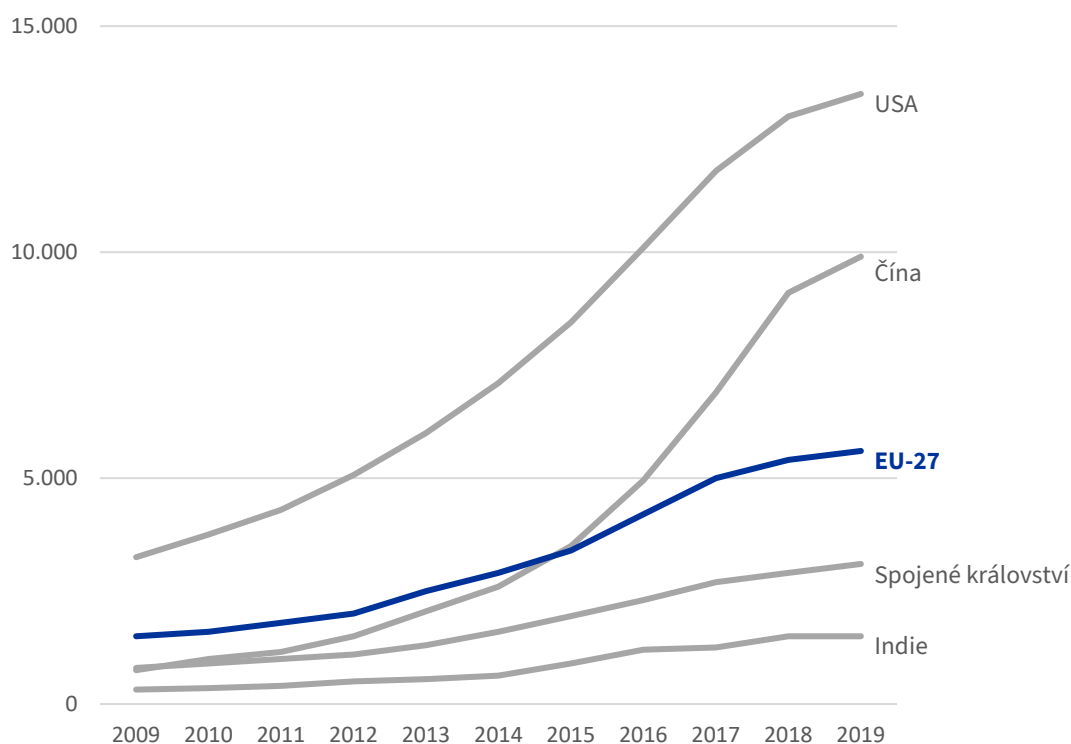
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

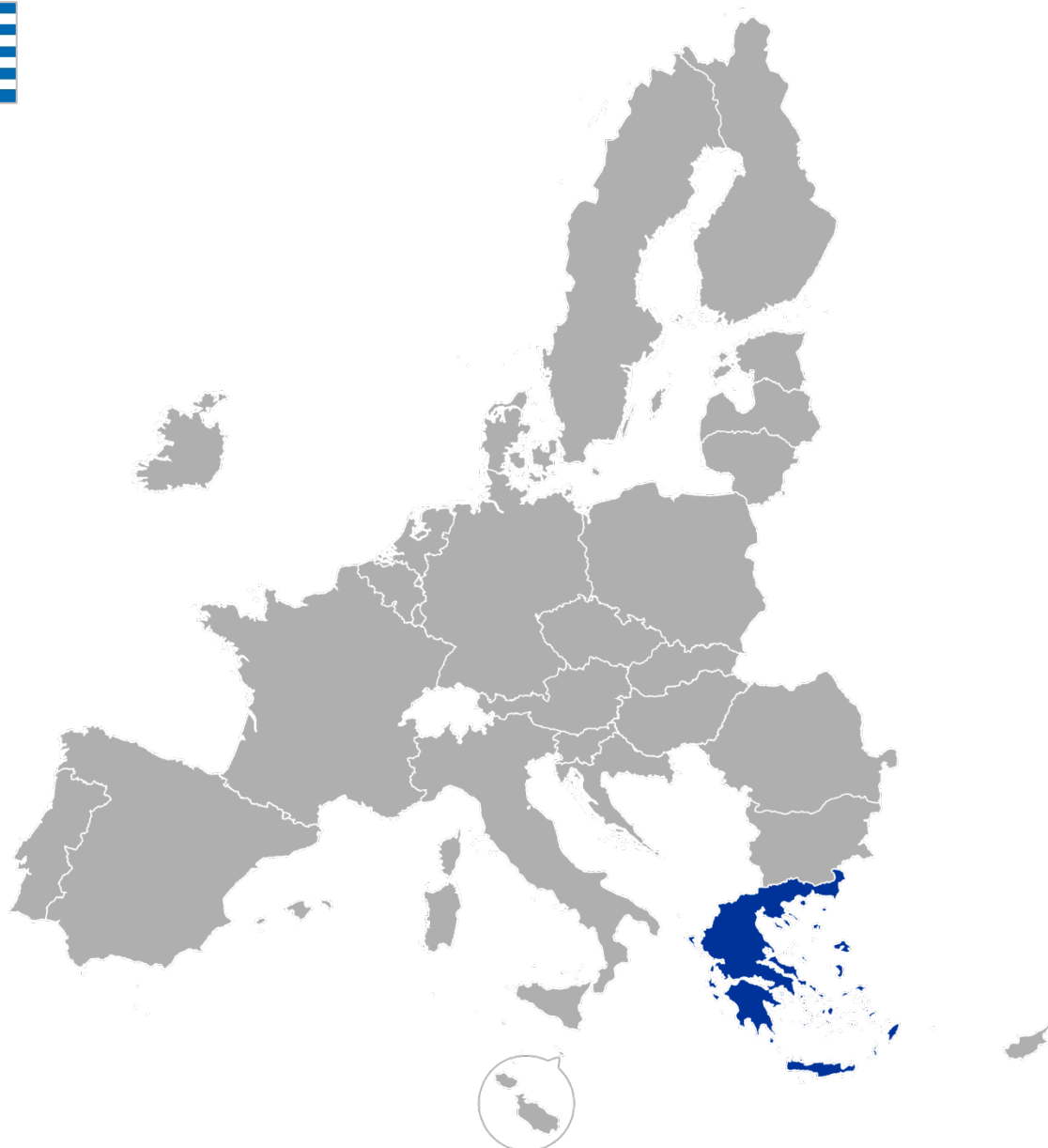
Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



ŘECKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



ŘECKO

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Belgie, Kypr, Německo a Španělsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uveďte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je pro přehlednost do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

Paní předsedkyně / pane předsedo, paní komisařko / pane komisaři, kolegyně a kolegové,

nejprve bych chtěl/a poděkovat Komisi za vypracování návrhu a předsednictví za jeho zařazení na pořad dnešního jednání.

Před dnešním zasedáním jsme vedli plodné diskuse s _____ .

Domníváme se, že je nezbytné přijmout klasifikaci rizik založenou na schopnostech/účelu, protože _____ .

Pokud jde o článek 1B, domníváme se, že během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat flexibilní přístup / přístup předběžné opatrnosti, protože _____ .

Pokud jde o článek 2, některé výjimky budou/nebudou platit, protože jsme pro _____ .

Věříme, že se nám dnes podaří nalézt přijatelný kompromis.

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vaše země se snaží vybudovat technologicky vyspělou budoucnost, která podporuje inovace a rozvoj ve prospěch všech a která se opírá o základní evropské hodnoty. Aby AI neměla negativní dopad na demokratické procesy, podporujete přístup založený na schopnostech, který považuje jakýkoli systém AI za vysoce rizikový, pokud tento systém dokáže „ovlivnit lidské chování“, například v kontextu voleb.
- Důrazně podporujete přístup předběžné opatrnosti ve vztahu k testování. Za nejdůležitější prvek považujete regulační pískoviště, která vývojáře nutí testovat systémy AI s malými skupinami uživatelů pod dohledem vnitrostátních orgánů. To zajišťuje bezpečné experimentování a snižuje riziko újmy.
- Zejména vás znepokojují potenciální rizika neobjektivní umělé inteligence: studie prokázaly, že chybovost nejpoužívanějších systémů pro rozpoznávání obličejů je mnohem vyšší při identifikaci osob s tmavší barvou pleti a žen. Nejvyšší chybovost je pak při kombinaci obou faktorů, tj. u žen tmavé pleti. Bylo zjištěno, že tyto chyby byly způsobeny trénovacími datovými soubory, které obsahovaly převážně tváře bílých mužů, což často vedlo k chybné identifikaci.
- Dosáhnout dnes společné shody může být náročné. **V případě potřeby by mohlo být jako kompromis navrženo dodatečné ustanovení o přezkumu po třech letech.** Tento přezkum by měl zahrnovat názory občanské společnosti, soukromého sektoru, příslušných odvětví a vlád členských států.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- I nadále je pro vás zásadní, aby se nařízení o AI nevztahovalo na záležitosti týkající se národní bezpečnosti, armády nebo obrany. Členské státy jsou odhodlány zachovat si kontrolu nad rozhodnutími v těchto oblastech, které stejně za normálních okolností nespádají do pravomoci EU. Řecko jakožto hraniční země schengenského prostoru zastává klíčovou úlohu při správě vnějších hranic EU. Systémy, které jsou v současné době zavedeny, jsou však zastaralé a neefektivní, což pro vnitrostátní orgány představuje velkou zátěž z hlediska provozu a humanitárních potřeb.
- Řecko proto uvažuje o využívání modelů AI k modernizaci a zefektivnění operací na hranicích – ke zlepšení procesů identifikace, registrace a správy případů a zároveň k podpoře rychlejšího a přesnějšího rozhodování. Vzhledem k naléhavé potřebě zvládnout vysoký počet vstupů na řecké území potřebuje vaše země flexibilní a včasné zavedení systémů AI na podporu rychlejších a efektivnějších operací.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

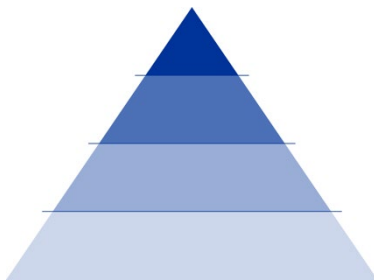
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko	na základě schopností	předběžná opatrnost	přezkum za 3 roky	armáda/obrana + národní bezpečnost	
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

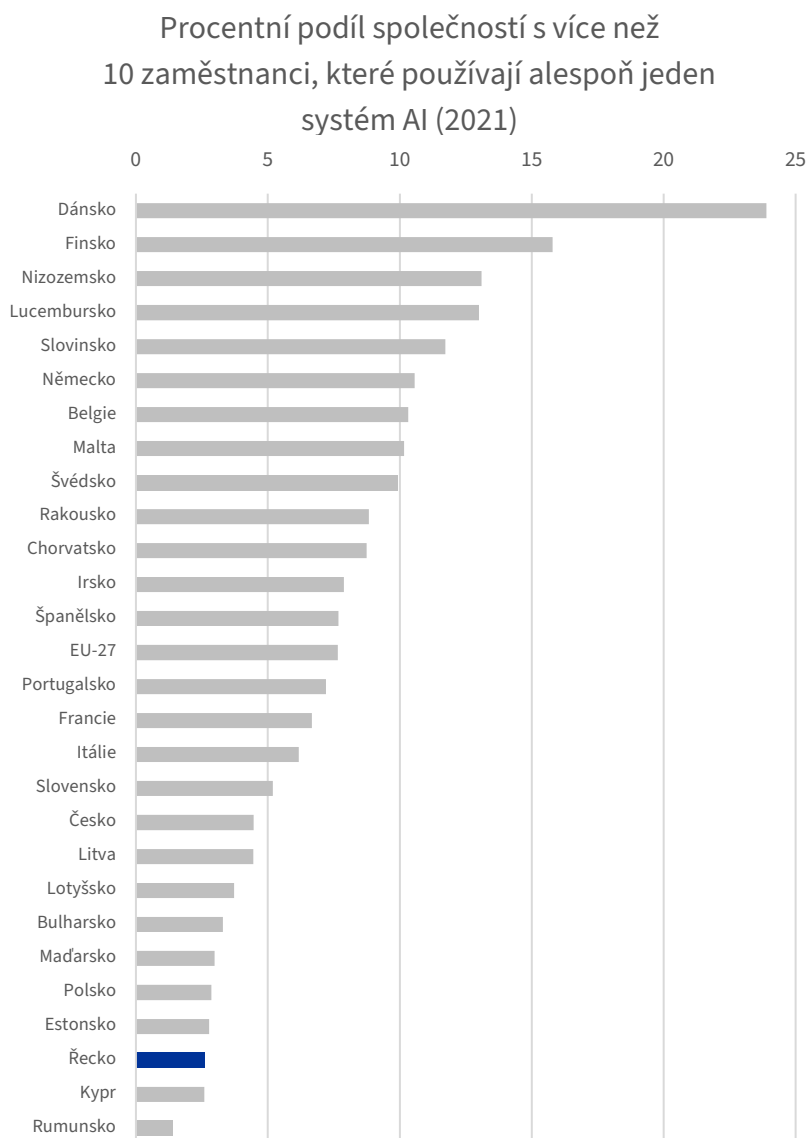
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

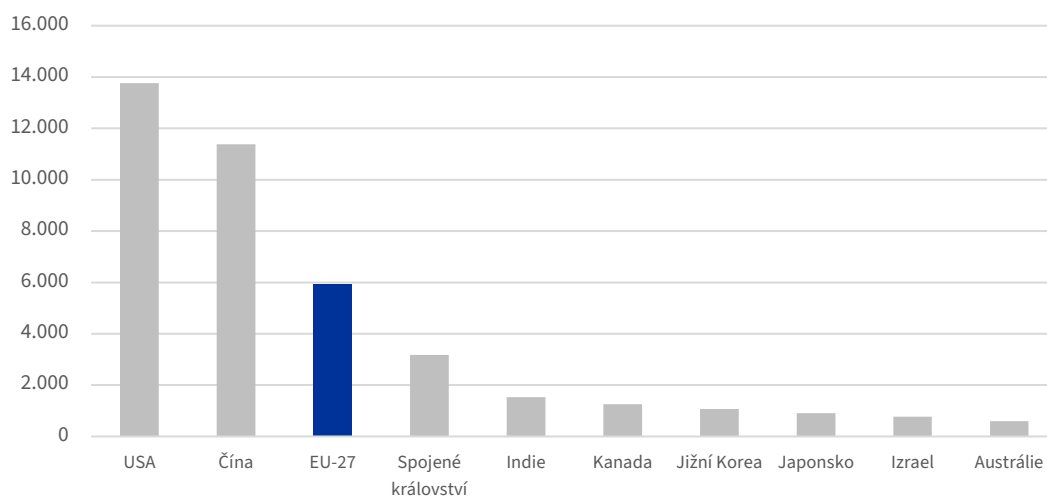
	Evropská unie	Řecko
Počet obyvatel	447 695 350	10 413 982
HDP na obyvatele (v EUR)	38 150	21 350
Míra nezaměstnanosti	6,1 %	11,1 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	4 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	20 %



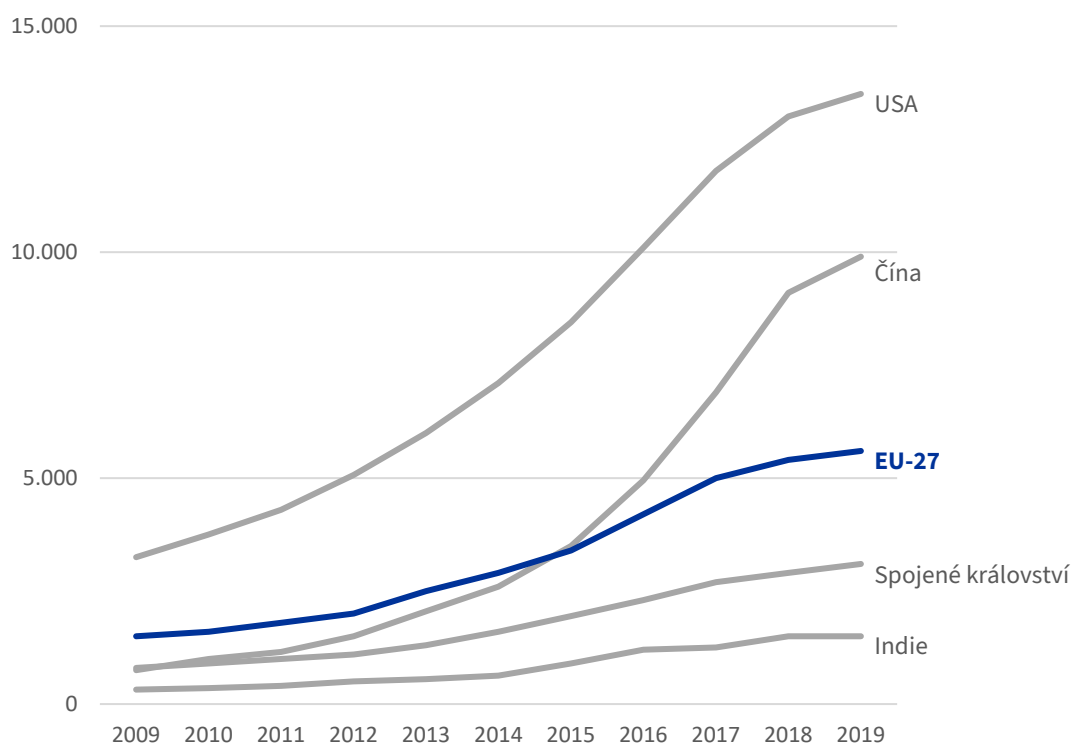
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



ŠPANĚLSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílají země, jako je Belgie, Kypr, Německo a Řecko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Prioritou vaší země jsou lidská práva a občanské svobody jako základ spravedlivé a bezpečné společnosti. Jste pro vnitrostátní investice do AI vyvíjené podniky a univerzitami, ale také považujete za zásadní, aby byla i nadále věnována pozornost etickým otázkám. Dnes je vaším hlavním cílem vytvořit rámec, který eliminuje i sebemenší riziko újmy způsobené umělou inteligencí, jako jsou zavádějící informace. Víte, že je to ambiciózní cíl, ale považujete jej za nezbytný.
- Jste zásadně proti návrhu Komise na zavedení přístupu založeného na účelu. Například pokud systém AI navržený pro kreativní odvětví k tvorbě hudby produkuje také nebezpečně hlasité zvuky (podobně jako zvuková zbraň), mohl by přece jen způsobit újmu, i když je označen za systém s nízkým rizikem. Soustředit se pouze na oznámený záměr a ignorovat technické schopnosti je nezodpovědné a nebezpečné.
- Pokud jde o povinnosti vývojářů, domníváte se, že nejzodpovědnějším krokem by bylo zavést přísné požadavky na testování, neboť AI je stále ještě v počáteční fázi vývoje. Pokud se vývojáři časem poučí ze svých chyb, podniky si budou více uvědomovat rizika spojená s předpojatostí a technologie se zdokonalí, jste ochotni jednat o možných výjimkách. Domníváte se však, že prozatím je nezbytné postupovat obezřetně a zavést silné záruky.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Vaše vláda je odhodlána prosazovat občanské svobody. Zároveň podporujete dvojí přístup k regulaci AI, který zajišťuje silnou ochranu v civilním prostoru a zároveň umožňuje flexibilnější přístup v oblasti obrany a národní bezpečnosti. Národní bezpečnost má zásadní význam pro zachování svobod, na nichž jsou založeny občanské svobody, a AI může hrát klíčovou úlohu při ochraně společnosti. Z těchto důvodů prosazujete, aby se nařízení o AI nevztahovalo na oblast armády, obrany a národní bezpečnosti.
- Pro tato odvětví by však měl být vytvořen samostatný regulační přístup vhodný pro daný účel, aby se zajistilo, že AI bude vyvíjena odpovědně. Požadavek plné transparentnosti by však mohl ohrozit utajení operací, zpozdit zavádění AI a snížit její účinnost v dynamickém nebo vysoce rizikovém prostředí.
- Případy jako španělský systém VioGén, který se používá k posuzování rizika domácího násilí, poukazují na nebezpečí, které představuje netransparentnost AI pod omezeným dohledem – situace některých obětí domácího násilí byla bohužel nesprávně klasifikována jako nízkoriziková a tyto osoby byly později zavražděny. Tyto případy selhání poukazují na nutnost odpovědného vývoje AI, zejména v oblasti bezpečnosti, ale neměly by být důvodem k nadměrné regulaci. EU by místo toho měla podporovat odpovědný přístup, zejména při vývoji AI pro bezpečnostní aplikace.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

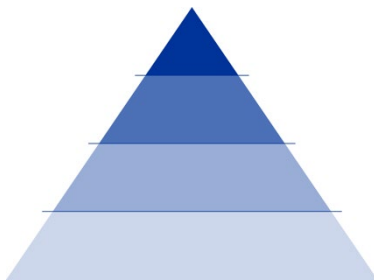
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabíjet bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko	na základě schopností	předběžná opatrnost		armáda/obrana + národní bezpečnost	
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

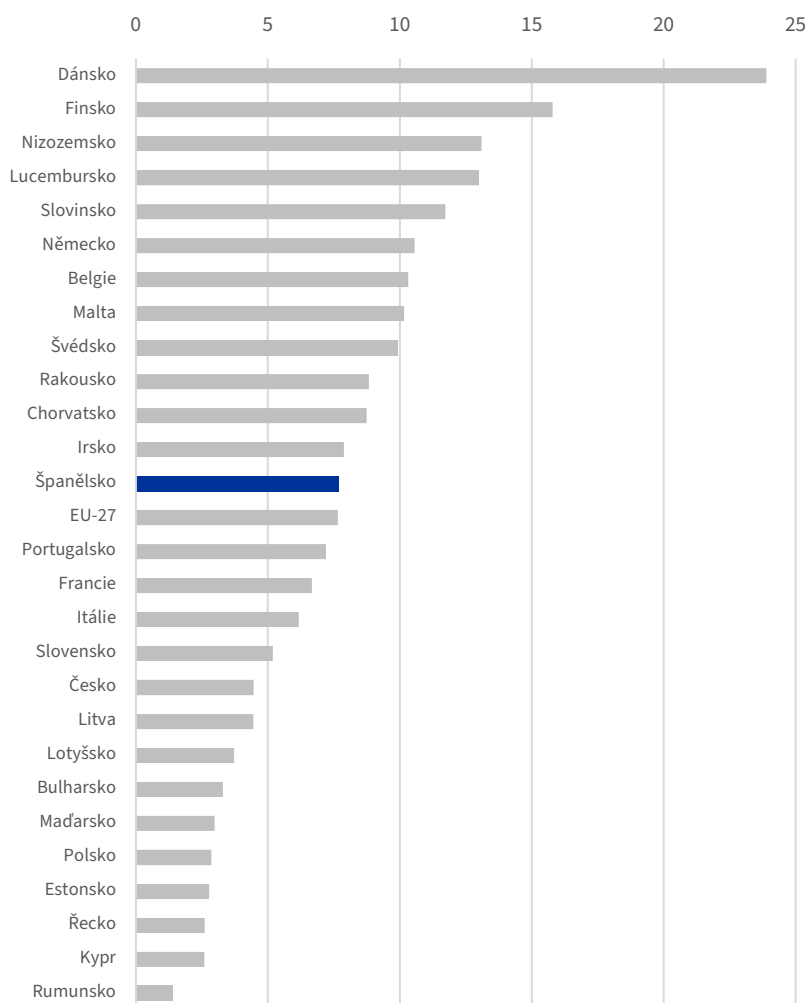
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Španělsko
Počet obyvatel	447 695 350	48 085 361
HDP na obyvatele (v EUR)	38 150	30 970
Míra nezaměstnanosti	6,1 %	12,2 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	9,2 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	38,7 %

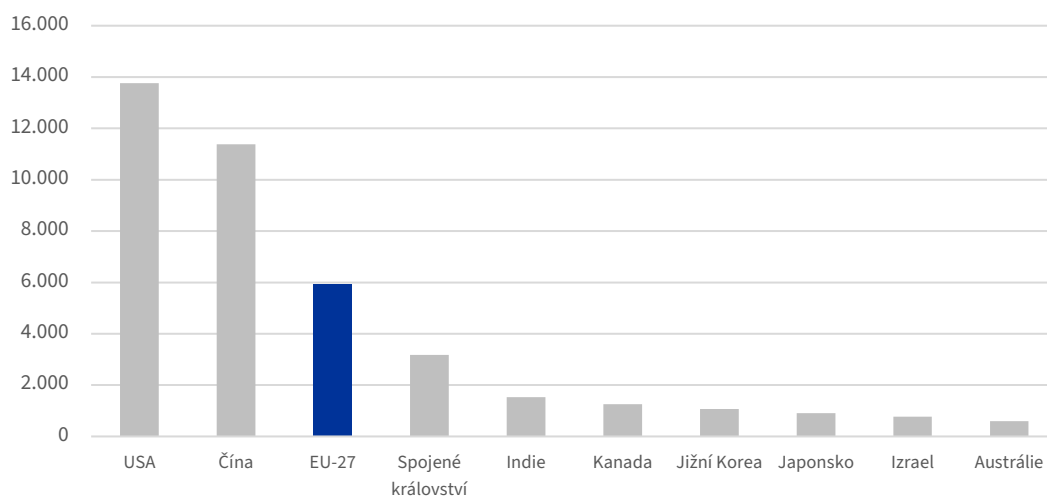
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



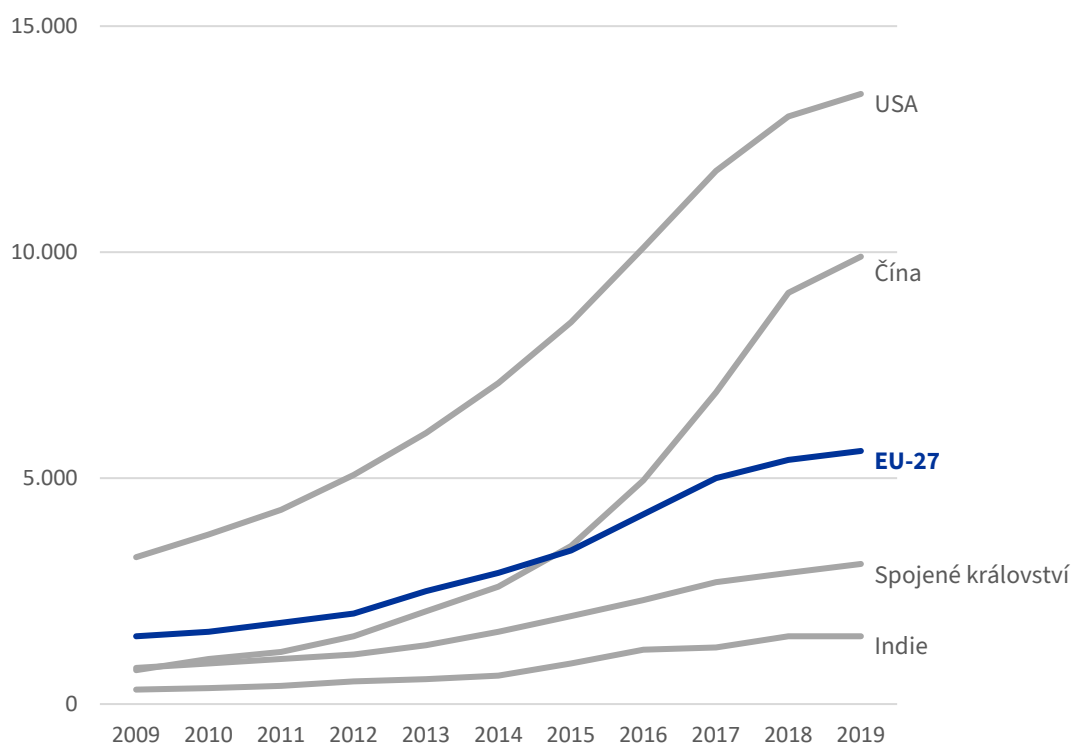
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

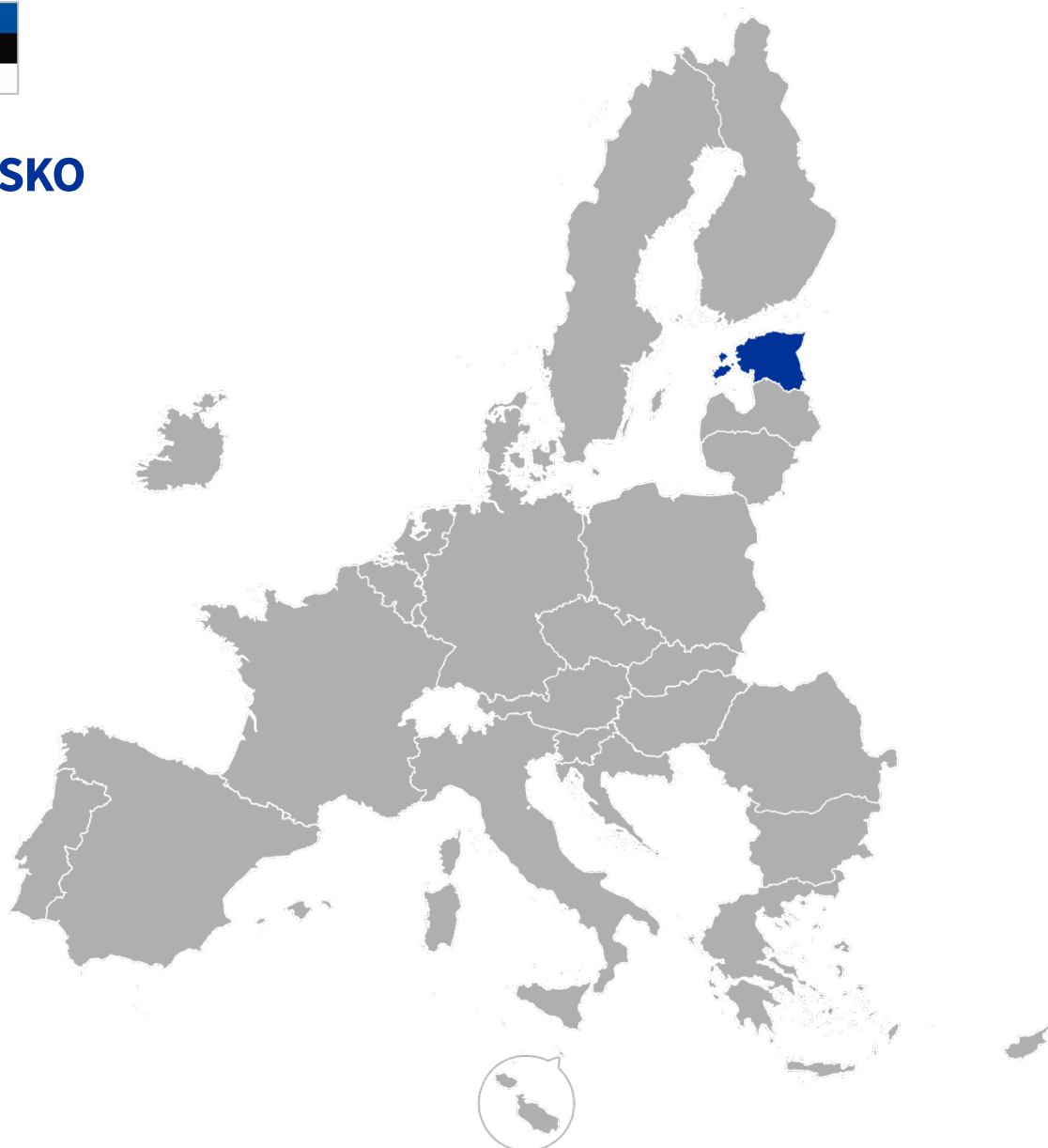
Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



ESTONSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Česko, Litva, Lotyšsko a Maďarsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vaše země podporuje jednotné nařízení EU o umělé inteligenci. Estonsko je jednou z nejvíce digitalizovaných zemí EU, a jeho obyvatelstvo je tedy velmi zběhlé v oblasti digitálních nástrojů, na rozdíl od některých zemí, kde si mnozí občané (zejména ti starší) s digitálními technologiemi tolik nerozumí a jsou zranitelnější vůči potenciálním rizikům.
- Návrh Komise však považujete za problematický. Nové systémy AI budou mít brzy mnohem větší schopnosti a hrozí, že definice „vysokého rizika“ vycházející pouze z aktuálních technických vlastností bude brzy zastaralá. Omezení technologií na základě potenciálních schopností může také bránit inovacím. V odvětvích, jako je zdravotní péče, by to mohlo znamenat, že se některá cenná řešení AI vyloučí dříve, než budou vůbec zvážena. Upřednostňujete přístup založený na účelu, neboť poskytuje výzkumným pracovníkům a vývojářům větší svobodu v oblastech s nízkým rizikem a zároveň jistotu ohledně požadavků na testování a silnější záruky v oblastech, kde jsou rizika naopak vyšší.
- Pokud jde o testování, podporujete flexibilní model. Společnosti se přirozeně snaží vytvářet bezpečné produkty, takže nadměrná regulace je zbytečná. **Považujete za důležité prosazovat testování nových modelů AI v reálných podmínkách**, neboť prostředí testovacího pískoviště často reálným podmínkám plně neodpovídá.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- I nadále je pro vás zásadní, aby do oblasti působnosti nařízení o AI nespadaly záležitosti týkající se národní bezpečnosti, armády nebo obrany. Členské státy jsou odhodlány zachovat si kontrolu nad rozhodnutími v těchto oblastech, které stejně za normálních okolností nespadají do pravomoci EU. To se odráží i ve skutečnosti, že Unie nemá jednotnou armádu EU.
- Dále argumentujete tím, že systémy AI používané pro shromažďování zpravodajských informací, kybernetickou obranu nebo plány vojenských operací nemohou podléhat stejným požadavkům na zveřejnění nebo regulačním požadavkům jako civilní systémy, protože by tak mohlo být ohroženo utajení operací, zájmy národní bezpečnosti a strategická výhoda, kterou mají tyto systémy poskytovat.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

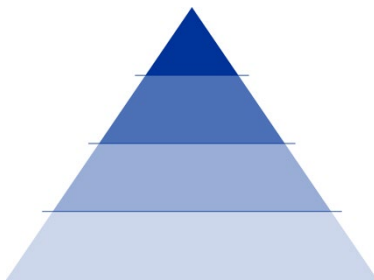
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko

Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko	na základě účelu	flexibilní	testování v reálných podmínkách	armáda/obrana + národní bezpečnost	
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

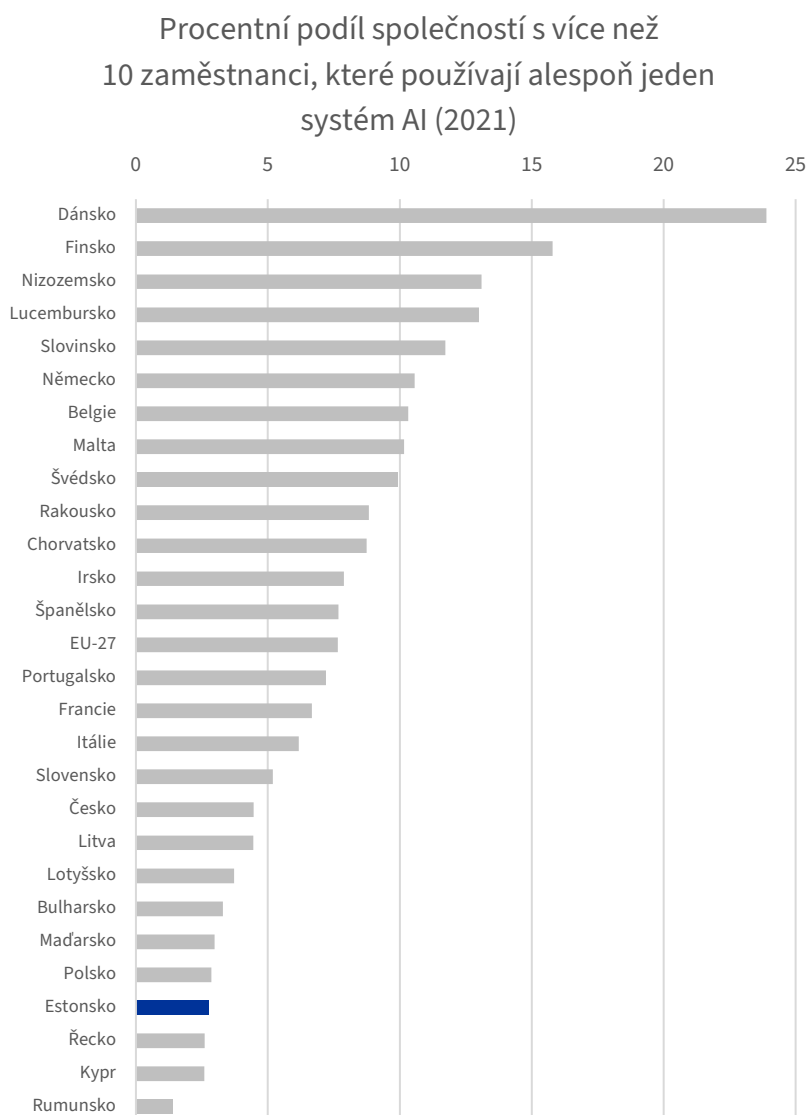
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

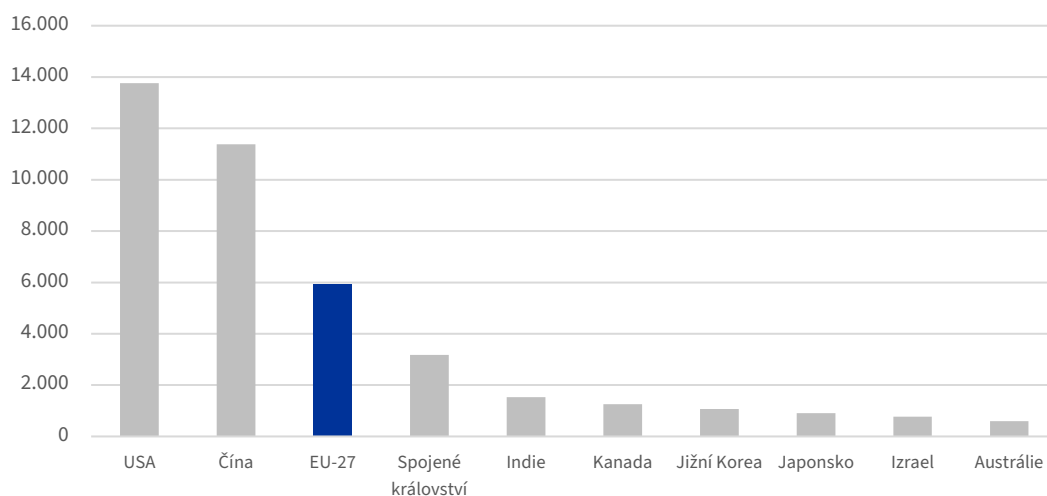
	Evropská unie	Estonsko
Počet obyvatel	447 695 350	1 365 884
HDP na obyvatele (v EUR)	38 150	27 960
Míra nezaměstnanosti	6,1 %	6,4 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	5,2 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	34,8 %



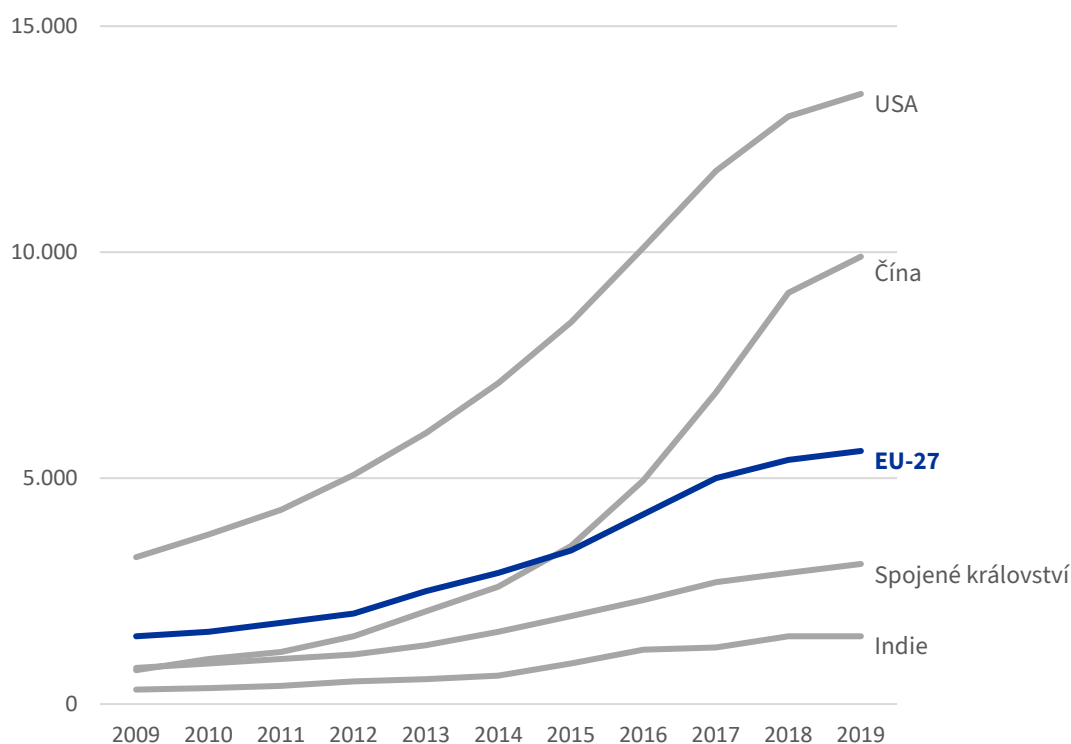
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



FINSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



FINSKO

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **prečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Irsko, Portugalsko a Švédsko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je pro přehlednost do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vaše země je jedním z lídrů EU v oblasti AI a jejím cílem je stát se průkopníkem digitální ekonomiky, která by byla spolehlivá, bezpečná a inovativní. V této oblasti máte také silnou podporu veřejnosti, částečně proto, že v oblasti AI existuje veřejné vzdělávání, které je občanům volně dostupné.
- Podporujete regulační přístup EU založený na účelu, který se zaměřuje spíše na to, jak je umělá inteligence využívána než na její technické vlastnosti. Předejde se tak tomu, aby AI byla označována za nebezpečnou, a bude možné soustředit se na její dopady v reálném světě, což považujete za nejlepší způsob, jak chránit občany.
- Návrh by podle vás měl klást důraz na podporu inovací ze strany EU prostřednictvím flexibilního přístupu, který mimo jiné umožňuje testování v reálných podmínkách pod dohledem veřejných orgánů. Je žádoucí členské státy podporovat v tom, aby zřizovaly regulační pískoviště, zavedení takové povinnosti by však mohlo odrazovat od zahraničních investic.
- Doufáte, že se dnes dohodnete na konečném znění aktu EU o umělé inteligenci, který bude vodítkem pro důvěryhodný vývoj AI v Evropě, ale zásadně nesouhlasíte s tím, aby byl v aktu stanoven přezkum po několika letech. Takové ustanovení by pouze narušilo regulační stabilitu a podkopalo důvěru investorů. Pokud by přetrvávala nejistota, **bylo by zodpovědnější odložit hlasování** než preventivně počítat se změnou pravidel.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- EU zatím v celosvětovém vývoji AI zaostává. Udělit výjimku startupům by mohlo pomoci formovat dynamický ekosystém. Startupy nemají právní, finanční ani správní kapacitu, aby byly schopny dodržovat stejná pravidla jako velké technologické společnosti. Přísné předpisy by mohly zbrzdit inovace v rané fázi.
- Pokud se vám nepodaří přesvědčit dostatečný počet členských států, aby přijaly flexibilní požadavky na testování uvedené v článku 1, chcete zajistit, aby startupům byly uděleny určité výjimky pro uplatňování těchto požadavků. Tímto způsobem nebudou přímo čelit úplné zátěži spojené s dodržováním předpisů (pokud to vyplývá z článku 1), ale budou moci fungovat v experimentálnějších prostředích.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

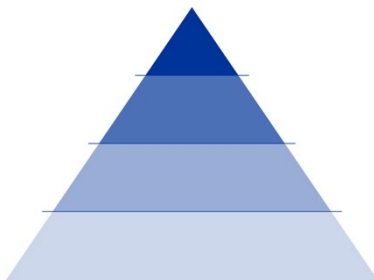
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabíjet bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko	na základě účelu	flexibilní	odložit hlasování	startupy	
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

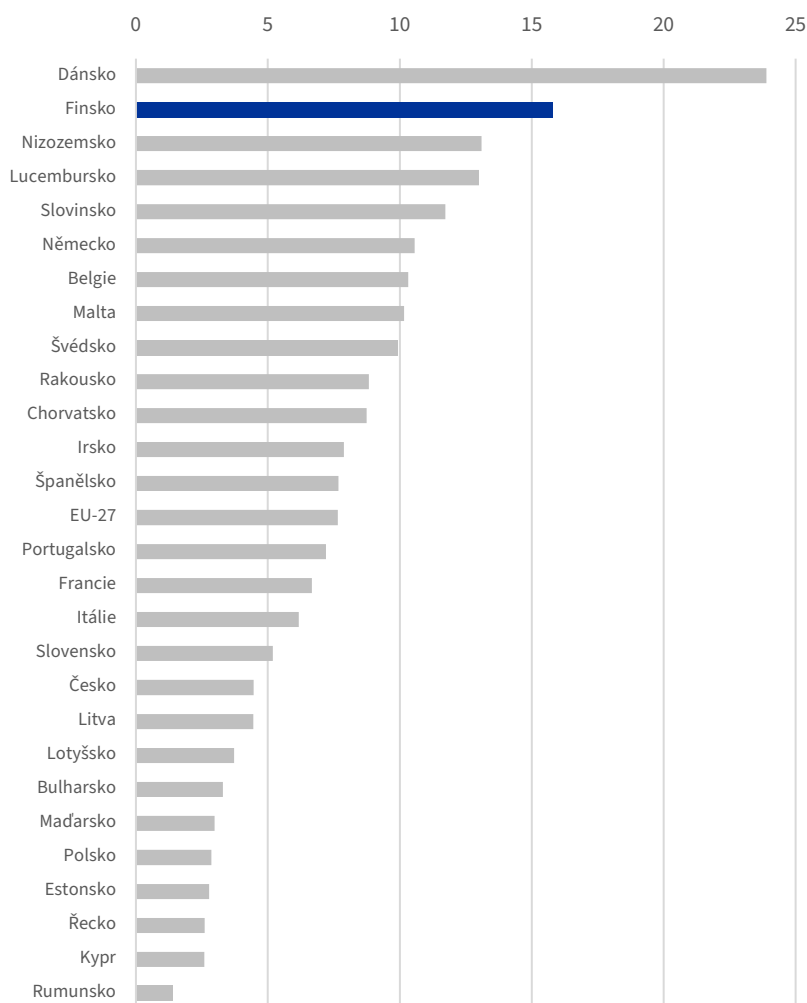
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Finsko
Počet obyvatel	447 695 350	5 563 970
HDP na obyvatele (v EUR)	38 150	48 910
Míra nezaměstnanosti	6,1 %	7,2 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	15,1 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	53,6 %

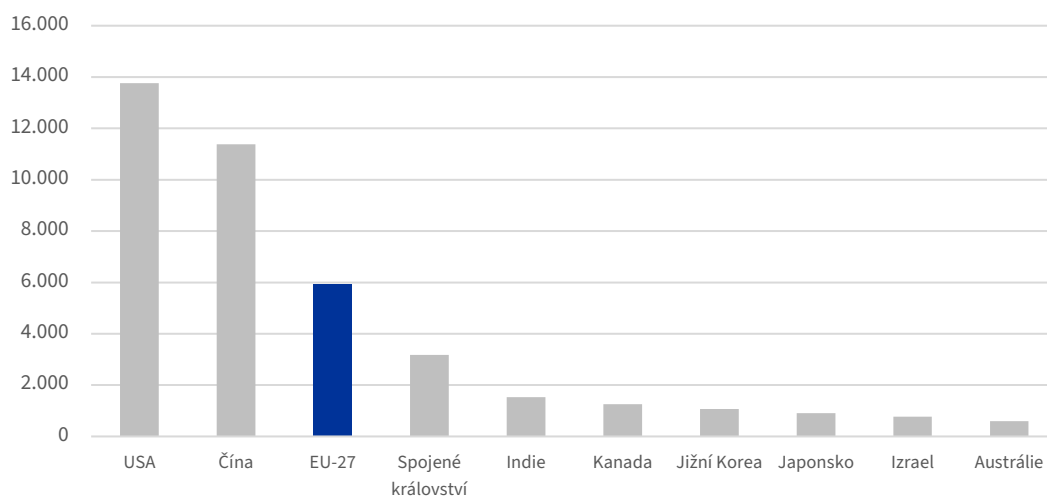
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



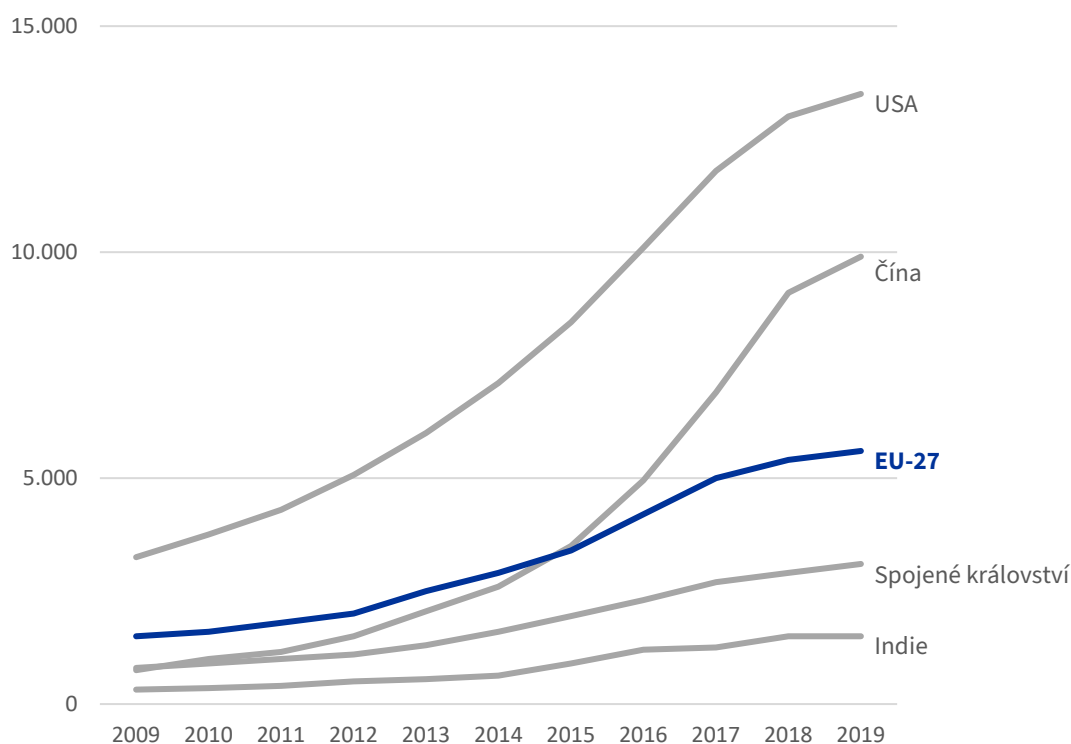
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

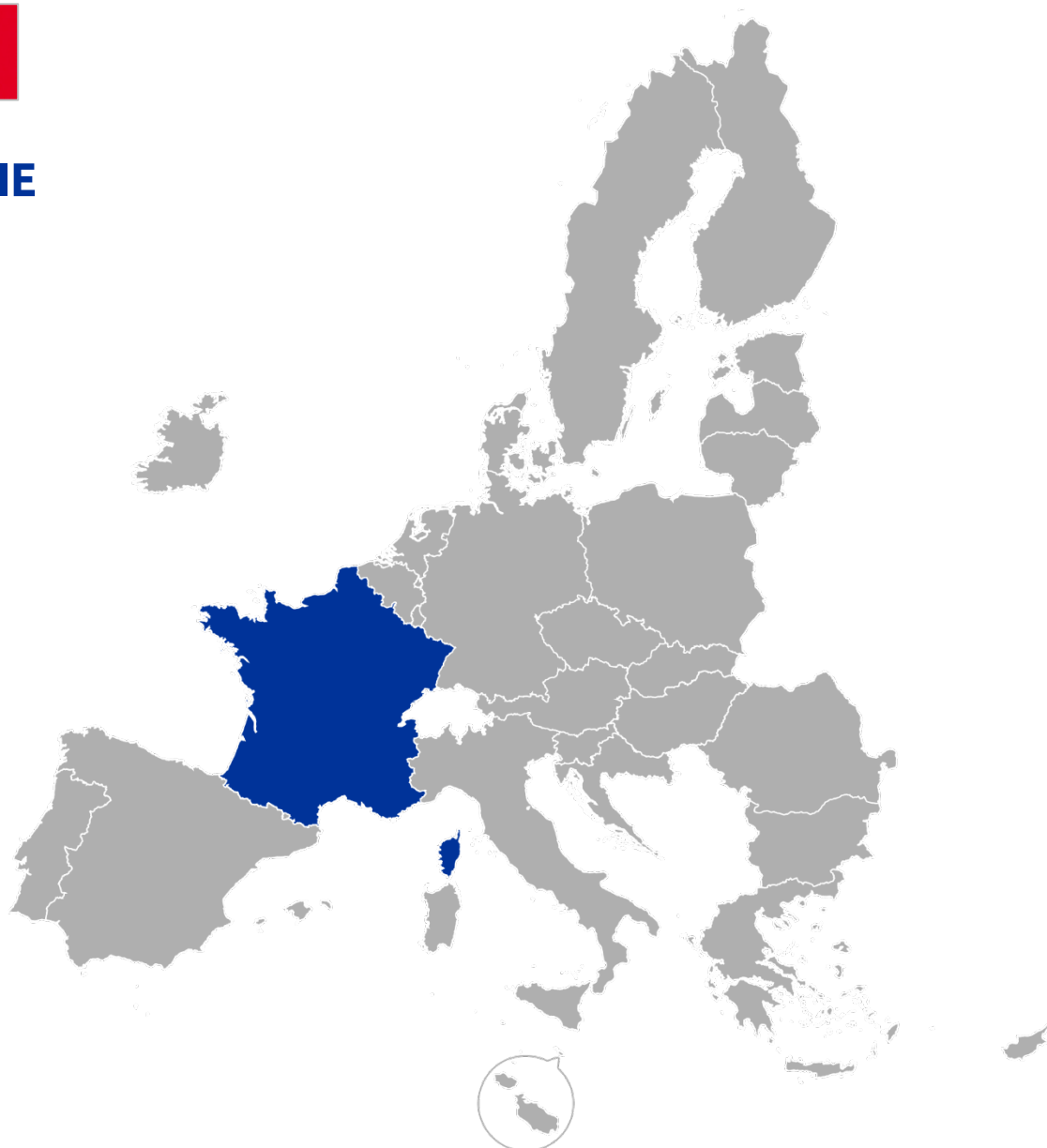
Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



FRANCIE



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



FRANCIE

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílají země, jako je Dánsko, Polsko, Rakousko a Rumunsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vláda vaší země již má strategii pro AI nazvanou „AI pro lidstvo“, která odráží vaši klíčovou prioritu: rozvoj zaměřený na člověka.
- Návrh Evropské komise přichází včas. V celosvětové diskusi o AI je naléhavě zapotřebí evropský hlas, který podpoří etický a odpovědný rozvoj. AI se dosud vyvíjela do značné míry bez kontroly, a i když existují rizika, neměla by být demonizována. Je i nadále silnou hnací silou inovací a růstu.
- Ačkoli podporujete přístup založený na účelu, protože je pro investory srozumitelný, **seznam citlivých odvětví by měl být rozšířen tak, aby zahrnoval i odvětví vzdělávání.** Chyby způsobené AI při rozřazování nebo hodnocení žáků a studentů mohou vést k diskriminaci a prohloubit nerovnosti. Spravedlivý přístup ke vzdělání má zásadní význam v boji proti strukturální diskriminaci.
- Ze zkušenosti víte, že pomáhají regulační pobídky – bez veřejného financování však hrozí, že inovativní nápady nebudou realizovány, což povede k odlivu mozků a migraci podniků. EU musí nejen přijmout rámec, ale také přislíbit významné veřejné investice do vývoje AI.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Vaše vláda je odhodlána prosazovat občanské svobody. Zároveň podporujete dvojí přístup k regulaci AI, který zajišťuje silnou ochranu v civilním prostoru a zároveň umožňuje pružnější normy v oblasti obrany a národní bezpečnosti. Národní bezpečnost má zásadní význam pro zachování základních svobod, na nichž jsou občanské svobody založeny, a AI může hrát při ochraně demokratických společností před vznikajícími hrozbami klíčovou úlohu.
- Proto prosazujete, aby se toto nařízení nevztahovalo na vojenské a obranné aplikace a aplikace v oblasti národní bezpečnosti a aby bylo vypracováno další nařízení o uplatňování AI v těchto odvětvích. Ačkoli je i nadále zásadně důležité, aby byla AI vyvíjena odpovědně, požadavek plné transparentnosti AI související s obranou by mohl ohrozit utajení operací, zpozdit její zavádění a snížit její účinnost v rychle se měnícím nebo nepřátelském prostředí.
- Kromě toho se zasazujete o to, aby bylo jasně stanoveno, že přísné požadavky na testování všech systémů AI se vztahují i na umělou inteligenci dovážanou ze zemí mimo EU. Navrhujete proto změnit znění na „Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI zaváděné v EU...“.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

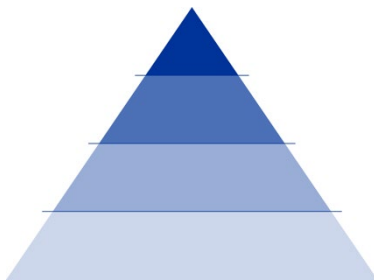
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie	na základě účelu	předběžná opatrnost	zahrnout oblast vzdělávání	armáda/obrana + národní bezpečnost	
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

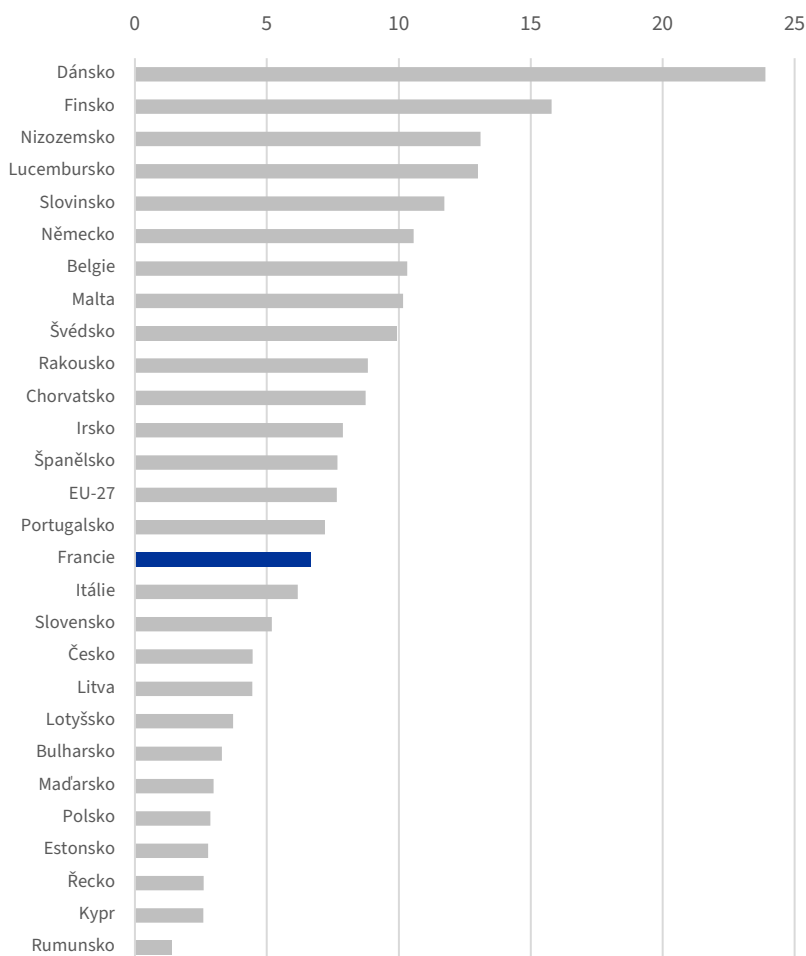
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Francie
Počet obyvatel	447 695 350	68 277 210
HDP na obyvatele (v EUR)	38 150	41 330
Míra nezaměstnanosti	6,1 %	7,3 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	5,9 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	30,3 %

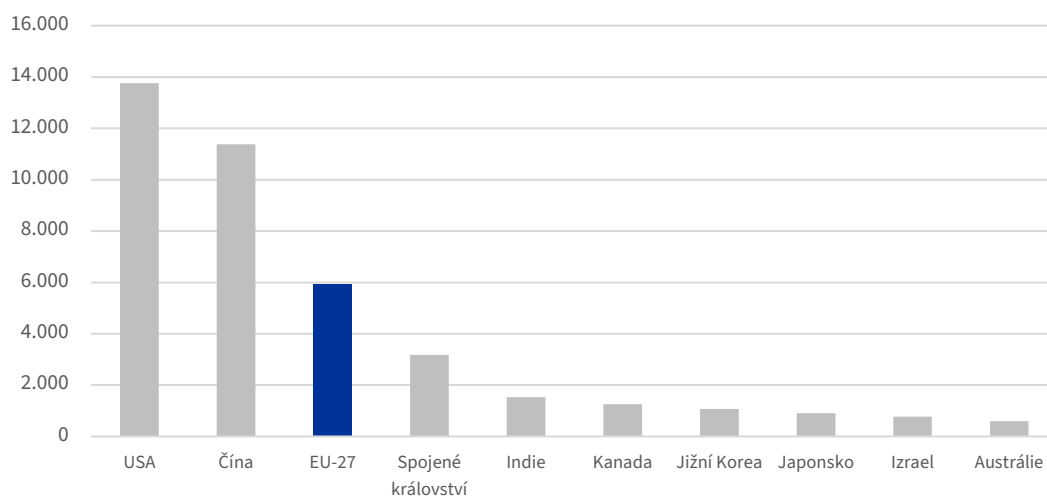
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



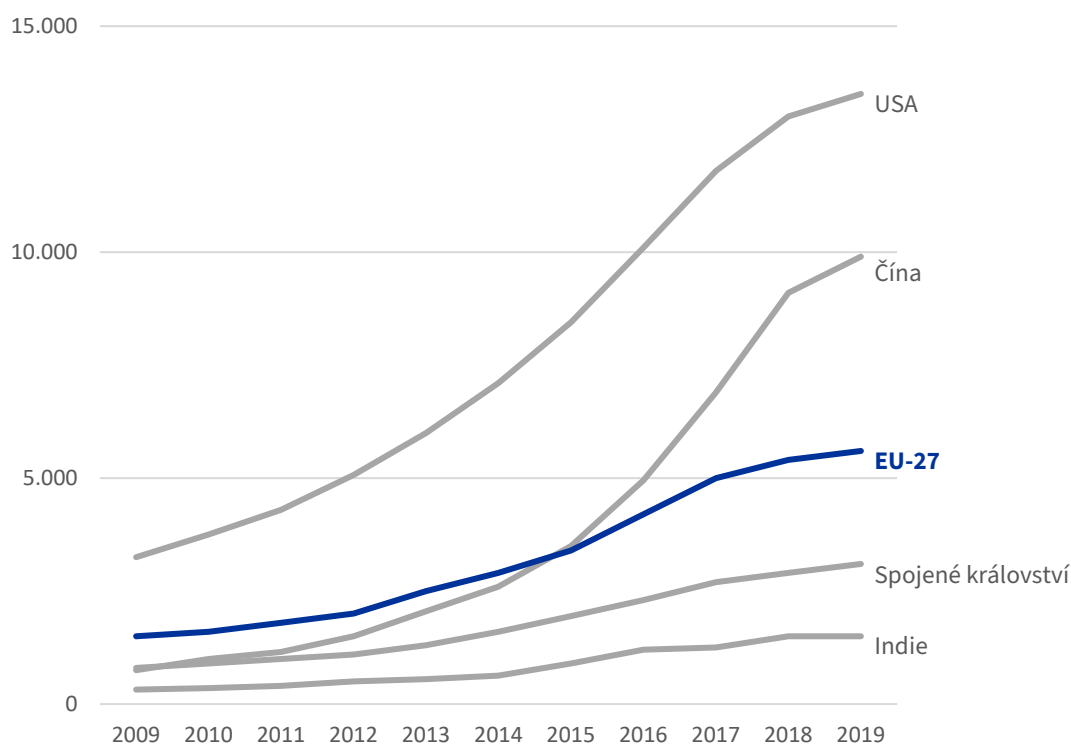
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



CHORVATSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Itálie, Nizozemsko a Slovinsko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost do **tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Počet startupů v Chorvatsku rychle roste. Podle odhadů poradenských společností by automatizace využívající AI mohla do roku 2025 generovat až 16 % HDP, a to zejména zvýšením produktivity. Automatizace by se mohla týkat až 52 % pracovní doby. Ačkoli to zní slibně pro hospodářský růst, vyvstávají zároveň obavy z potenciální nezaměstnanosti ve výrobním odvětví.
- Co se týče klasifikace umělé inteligence, podporujete přístup Komise založený na účelu, neboť reálně odráží specifická rizika – například rozpoznávání obličeje pro odemčení telefonu vs. rozpoznávání obličeje při hromadném sledování.
- Jste však přesvědčeni, že transparentnost je klíčová: všechny systémy AI bez ohledu na úroveň rizika by měly uživatele jasně informovat o svém účelu. To zajišťuje odpovědnost a chrání práva jednotlivců.
- Podporujete také přísné normy pro vývoj AI, nicméně stávající pokyny pro posuzování rizik považujete za příliš vágní. Zasazujete se o povinné posuzování dopadu vysoce rizikové AI na lidská práva. Postup takového posuzování by měl být vyvinut na základě konzultací s dotčenými komunitami, aby byly chráněny zranitelné skupiny.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Podle vašeho názoru by akt EU o umělé inteligenci měl urychlit rozvoj AI ve všech členských státech Unie. Zavedené společnosti působící v oblasti AI sídlí v několika zemích s velkým kapitálem. Oproti tomu jiné členské státy jsou stále ve fázi, kdy se snaží vytvářet příznivé podmínky pro startupy. Podstatou startupů je vytvořit v krátkém čase a s omezeným přístupem ke kapitálu něco převratného. To je při velké administrativní zátěži téměř nemožné.
- Proto byste chtěli, aby se na startupy (s méně než 10 zaměstnanci a ročním obratem nižším než 2 miliony eur) nařízení EU o umělé inteligenci nevztahovalo.
- Pokud se vám nepodaří přesvědčit dostatečný počet dalších členských států, aby váš postoj podpořily, považujete za dobrý kompromis to, aby EU vyhradila pro startupy dodatečné zdroje. Jejich činnost tak nebudou brzdit administrativní a rozpočtová omezení vyplývající z tohoto nařízení.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

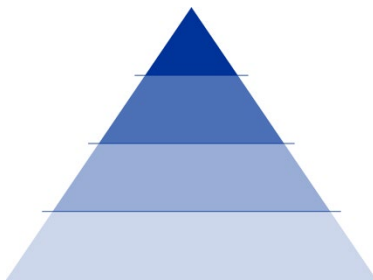
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko	na základě účelu	předběžná opatrnost		startupy	
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

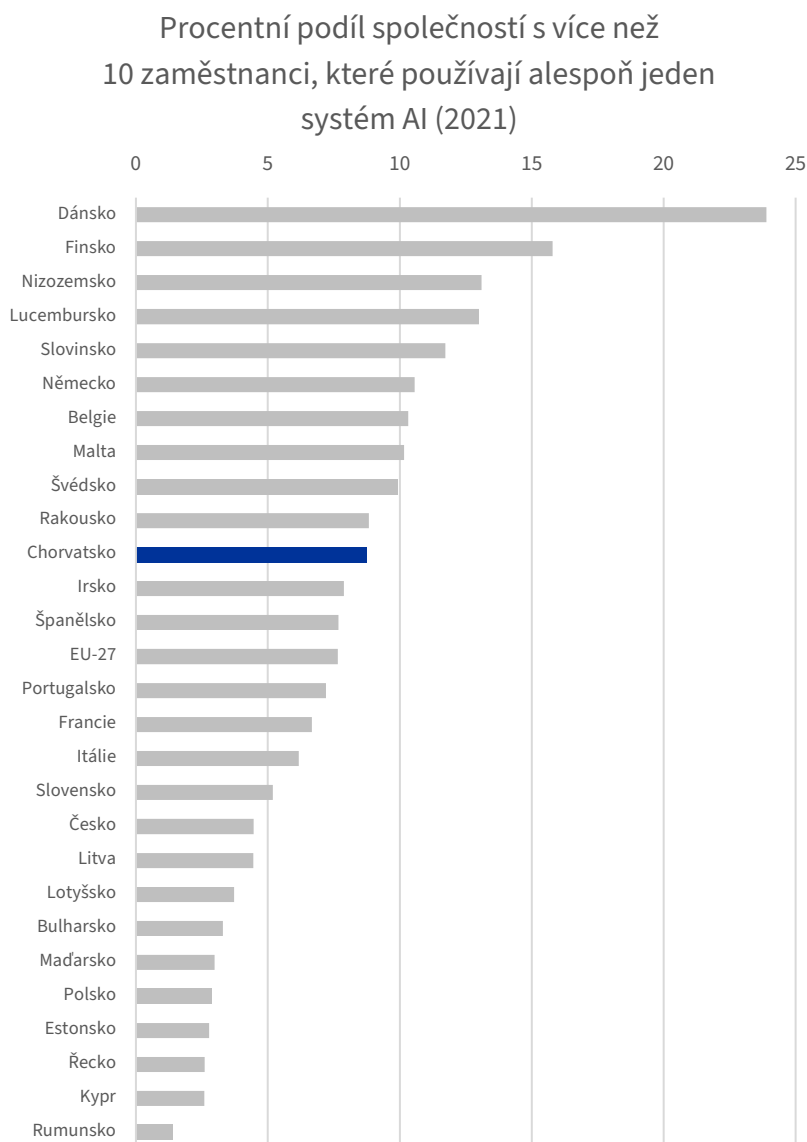
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

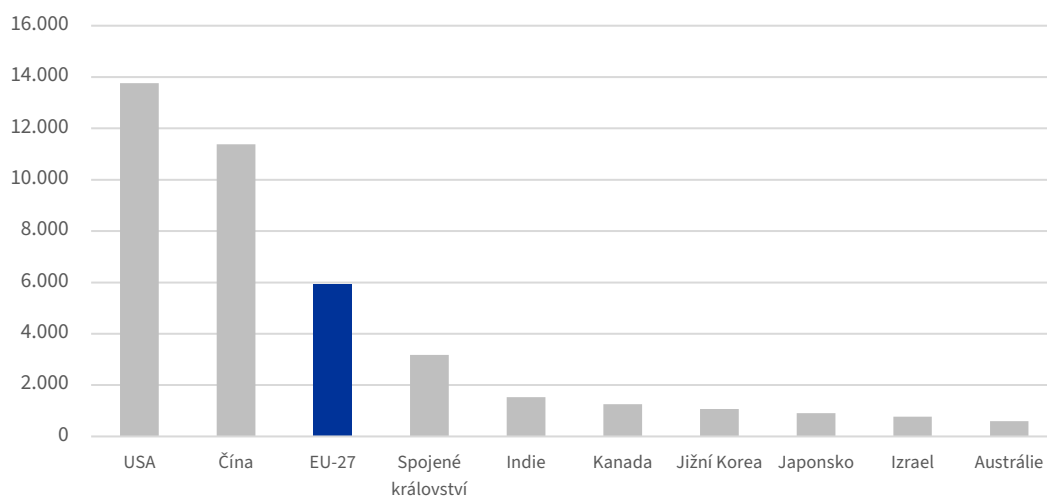
	Evropská unie	Chorvatsko
Počet obyvatel	447 695 350	3 850 894
HDP na obyvatele (v EUR)	38 150	19 800
Míra nezaměstnanosti	6,1 %	6,1 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	7,9 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	25 %



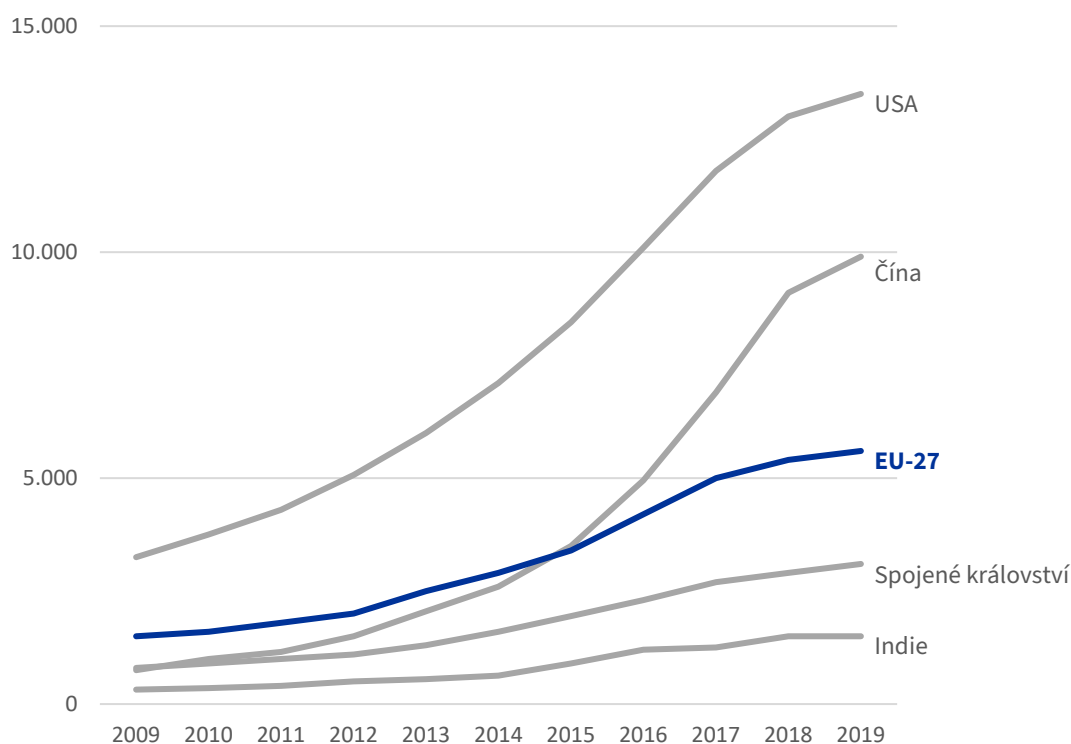
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



MAĎARSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



MAĎARSKO

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **prečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílejí země, jako je Česko, Estonsko, Litva a Lotyšsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je pro přehlednost do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vaše vláda nedávno vypracovala strategii pro AI. Jedním z jejích pilířů je připravit společnost na praktické zvládnání nevyhnutelných změn, které AI přináší, a umožnit, aby výhody této technologie mohly být využívány v průmyslové i akademické sféře.
- Stále rozsáhlejší využívání AI by podle vás mohlo přinést významnou přidanou hodnotu pro národní HDP – zdroje, které jsou nezbytně nutné pro investice do veřejné infrastruktury.
- Pokud jde o klasifikaci „vysoce rizikových“ modelů AI, považujete přístup založený na schopnostech vhodnější než přístup založený na účelu, který navrhuje Komise. Je-li stejný systém AI považován v jednom kontextu za nízkorizikový a v jiném za vysoce rizikový, může to pro uživatele, vývojáře a investory představovat nejasnosti a právní nejistotu.
- Máte za to, že Evropská unie se musí výrazně podílet na zajišťování spravedlivé hospodářské soutěže, aby evropské i maďarské společnosti zůstaly konkurenceschopné na vnitřním trhu a v globálních hodnotových řetězcích. Maďarsko se chce stát atraktivní destinací pro zahraniční investice. Proto podporujete regulaci, v níž budou stanoveny pouze minimální požadavky nebo normy, aby se zabránilo tomu, že vývojáři budou zatíženi vysokými náklady.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Zasazujete se o to, aby se regulace AI nevztahovala na vojenské a obranné aplikace. Také zdůrazňujete, že nadměrná regulace AI by mohla bránit rychlému zavádění technologií AI v oblastech zasažených konfliktem nebo během reakce na krizi. Region, ve kterém se Maďarsko nachází, se již dlouhou dobu potýká s nestabilitou, a i proto je při používání technologií AI zásadní flexibilita.
- Vzhledem k tomu, že Maďarsko sdílí hranici s Ukrajinou, která je v současné době ve válce s Ruskem, má závažné obavy ohledně své národní bezpečnosti a bezpečnosti svých občanů. Jako hraniční stát schengenského prostoru navíc čelí značným výzvám, pokud jde o řízení zvýšených migračních toků, zejména na svých jižních hranicích.
- S ohledem na tyto tlaky očekává Maďarsko větší pochopení a podporu ze strany ostatních členských států EU při zvládání svých specifických geopolitických a bezpečnostních výzev.
- I když jsou některé země znepokojeny riziky dovážených systémů AI, zejména z Číny, domníváte se, že mnohé z těchto obav jsou spíše politicky motivované, než že by byly založeny na skutečných hrozbách. Maďarsko má s čínskými technologickými partnerstvími dobré zkušenosti a podle vás by pro EU bylo přínosné, kdyby se napětí v oblasti obchodu s Čínou vyřešilo.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

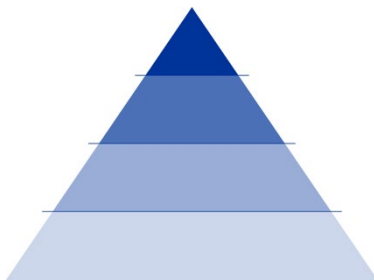
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko	na základě schopností	flexibilní		armáda/obrana	
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

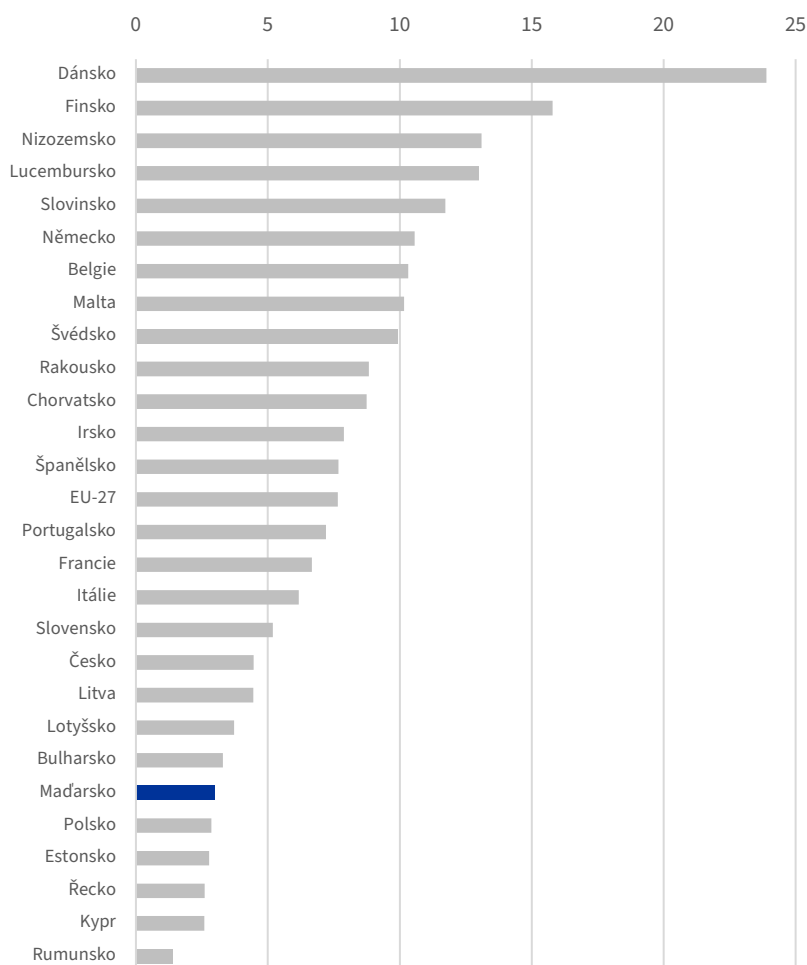
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Maďarsko
Počet obyvatel	447 695 350	9 599 744
HDP na obyvatele (v EUR)	38 150	20 630
Míra nezaměstnanosti	6,1 %	4,1 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	3,7 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	28,1 %

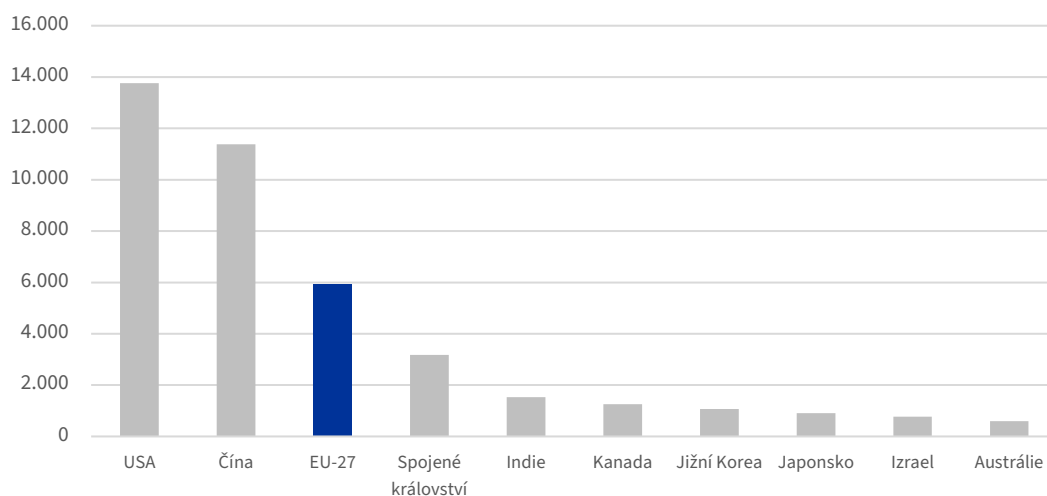
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



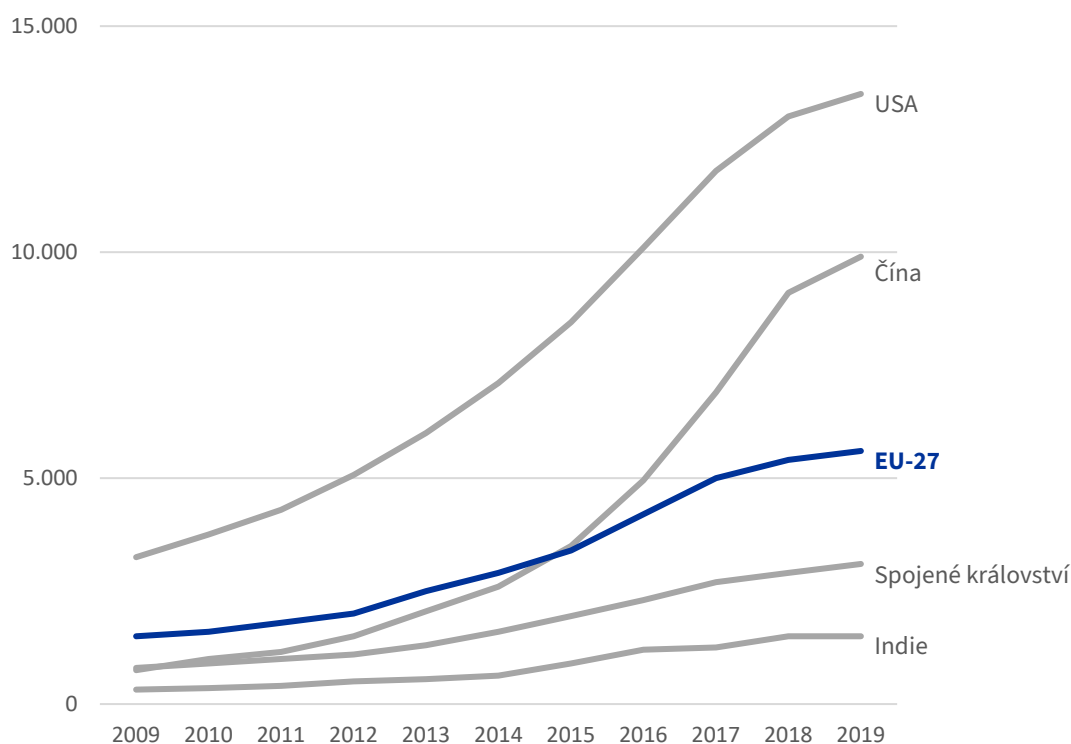
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

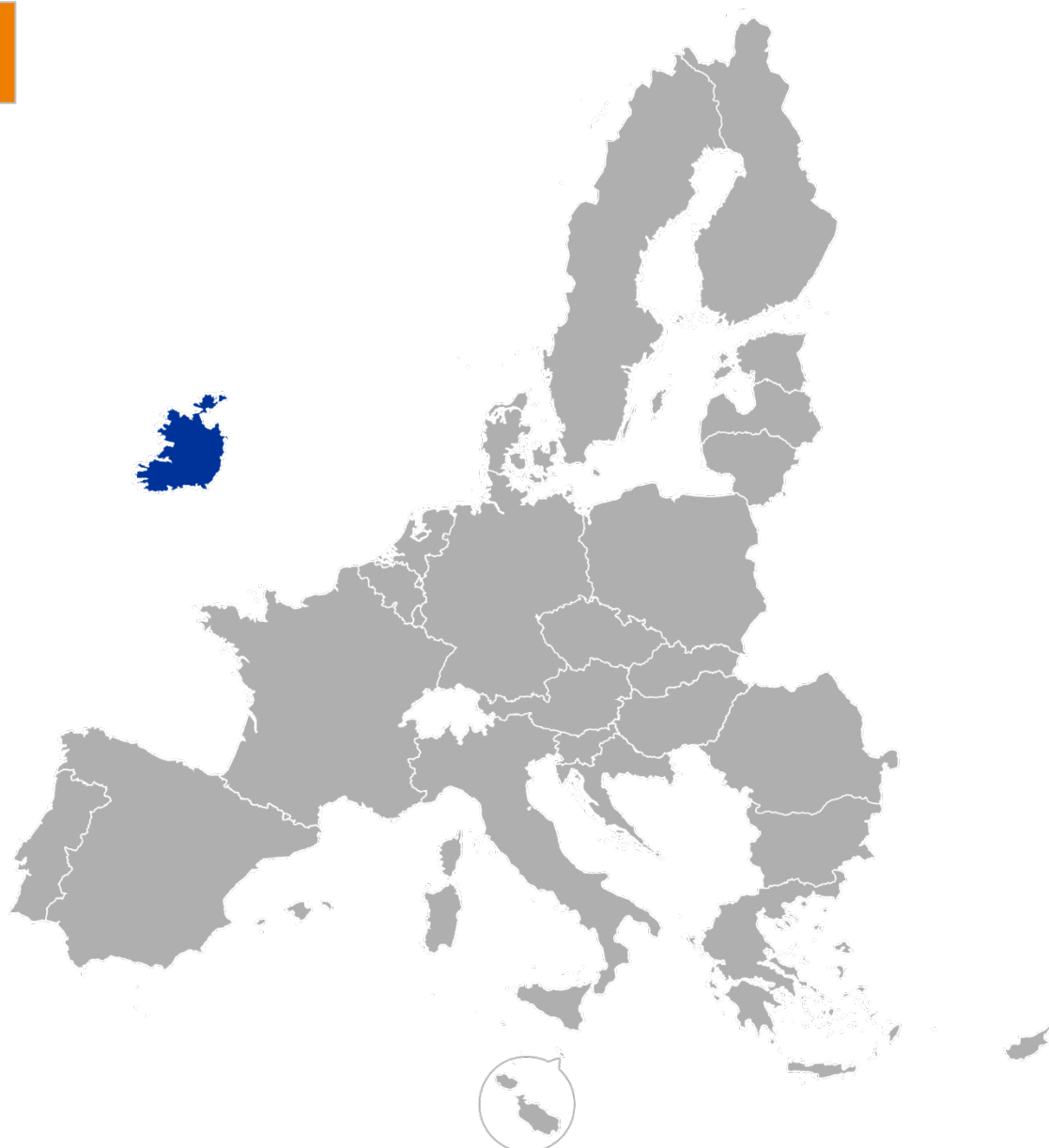
Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



IRSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



IRSKO

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílají země, jako je Finsko, Portugalsko a Švédsko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost do **tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Ve vaší zemi se nacházejí evropské centrály společností, jako je OpenAI, Google a Meta. Přítomnost těchto společností v Irsku je vaším národním zájmem, ale zároveň přináší výhody i pro Unii jako celek.
- Vítejte návrh Komise na zavedení společného rámce EU pro umělou inteligenci. Považujete ho za důležitý, protože vysílá okolnímu světu jasný signál, že je EU připravena stát se odpovědnou velmocí v oblasti AI. Podporujete proto návrh Komise na zavedení přístupu založeného na účelu. U tohoto přístupu je (na rozdíl od přístupu založeného na schopnostech) menší riziko, že vznikne rozpor s požadavky stanovenými v předpisech regulujících jiná odvětví.
- Pokud jde o požadavky na vývojáře, **zdůrazňujete potřebu flexibility – právní úprava nesmí odrazovat investory a musí zároveň zajistit odpovědnost.**
- Upozorňujete také na skutečnost, že mnoho rizik spojených s umělou inteligencí vzniká až po jejím zavedení. Je tudíž účinnější průběžně monitorovat než odsouvat datum uvedení na trh na základě hypotetických rizik.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- AI a digitální revoluce musí podle vás také podporovat genderovou rovnost v EU. V odvětvích souvisejících s výzkumem a inovacemi jsou ženy nedostatečně zastoupeny téměř ve všech oblastech a na všech úrovních. Zdůrazňujete, že je důležité aktivně podporovat zapojení žen do rozhodování o dalším formování AI.
- Startupy představují jedinečnou příležitost. Na rozdíl od velkých technologických společností existují ve startupech často menší genderové rozdíly. Mimo to i v oblasti AI nabízejí startupy více možností pro ženy a genderqueer osoby. Je také prokázáno, že různorodé týmy budují méně předpojatou a etičtější umělou inteligenci – něco, co by EU měla podporovat.
- Ze své podstaty usilují startupy o rychlé inovace s omezenými zdroji. Tento potenciál může potlačit velká administrativní zátěž, přičemž tomuto riziku jsou vystaveny zejména malé týmy, které se pokoušejí prosadit.
- Proto podporujete, aby malé startupy (s méně než 10 zaměstnanci a obratem do 2 milionů eur) byly vyňaty z plného rozsahu nařízení EU o umělé inteligenci.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

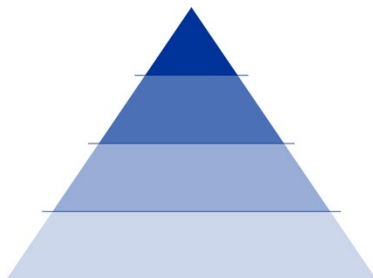
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabíjet bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko	na základě účelu	flexibilní	flexibilita při testování	startupy	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

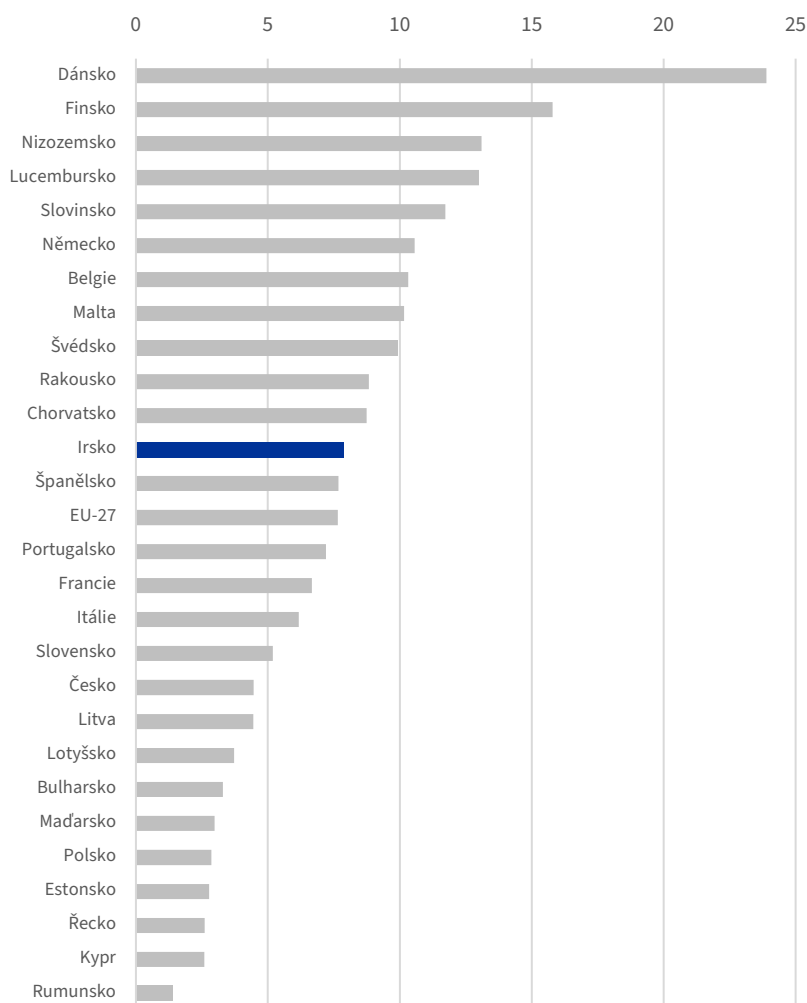
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Irsko
Počet obyvatel	447 695 350	5 271 395
HDP na obyvatele (v EUR)	38 150	96 290
Míra nezaměstnanosti	6,1 %	4,3 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	8 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	43,8 %

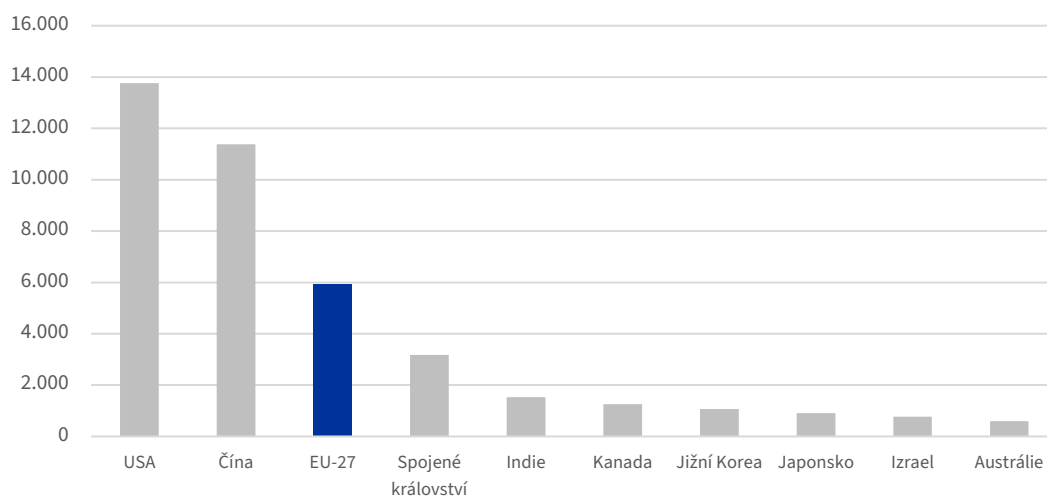
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



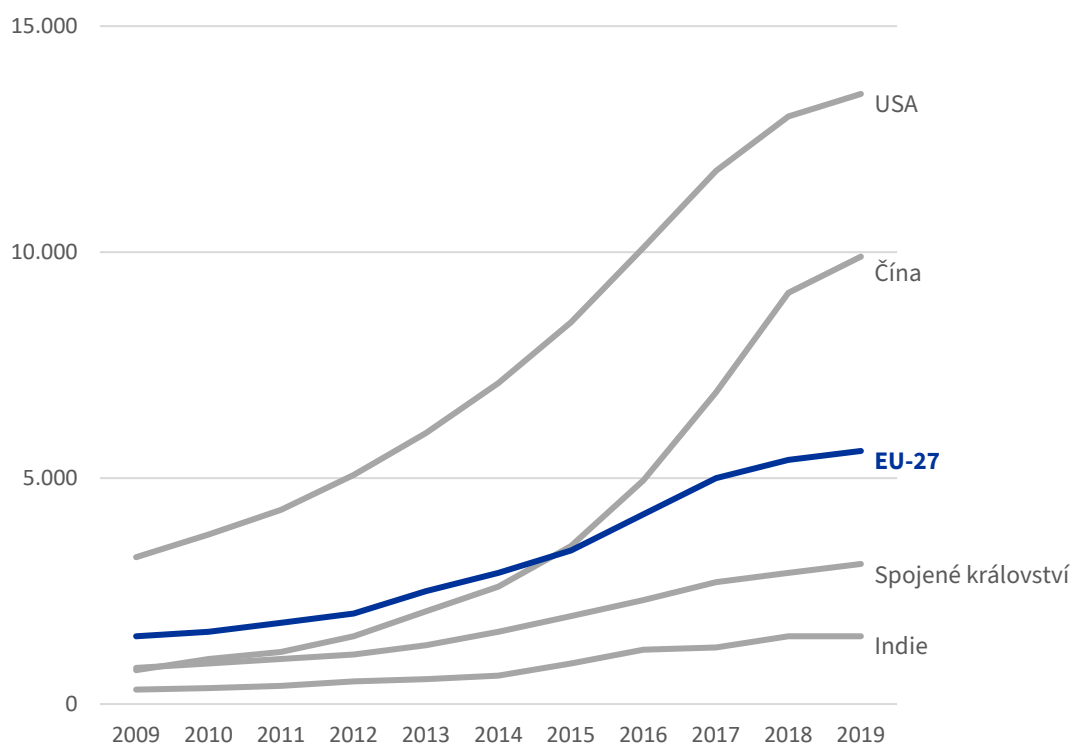
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



ITÁLIE



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **prečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Chorvatsko, Nizozemsko a Slovinsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože _____*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože _____*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro _____*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vláda vaší země se domnívá, že regulace AI by se měla řešit na evropské úrovni, protože stanovení pravidel jednotlivými členskými státy by mohlo oslabit konkurenceschopnost a strategickou autonomii EU jako celku.
- Na základě článku 3 italské ústavy („Všichni občané mají stejnou sociální důstojnost a jsou si rovni před zákonem.“) a v souladu s cíli udržitelného rozvoje OSN musí AI podporovat inkluzivnost, udržitelnost a lidská práva a být k užítku lidem i planetě.
- Přístup založený na účelu považujete za problematický, protože vede k rozporům v rámci jednotlivých odvětví. Například chirurgický výkon za pomoci AI se klasifikuje jako vysoce rizikový, stejně jako využití AI pro přiřazování nemocničních místností – navzdory tomu, že úroveň dopadu použití AI je v každé z těchto situací odlišná. Přístup založený na schopnostech podle vás lépe odráží nebezpečí, která hrozí v reálném životě. Podporujete však také pravidelnou aktualizaci seznamu rizik tak, aby se zohlednily vyvíjející se schopnosti AI.
- Upřednostňujete také zavedení preventivního testování – nejen kvůli obavám o bezpečnost, ale také z důvodů udržitelnosti. Vývoj umělé inteligence je totiž náročný z hlediska pracovní síly, ale i energie. Například pro trénování velkých jazykových modelů se spotřeba elektřiny pohybuje ve stovkách megawatthodin. Neregulovaný přístup typu pokus – omyl by tudíž zvyšoval náklady na ochranu životního prostředí. Přísnější právní rámec může mírně zpomalit vývoj, ale v dlouhodobém horizontu přispěje k tomu, aby byla AI bezpečnější, účinnější i důvěryhodnější.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Podle vašeho názoru by akt EU o umělé inteligenci měl urychlit rozvoj AI ve všech členských státech Unie. Startupy mají klíčovou úlohu při podpoře inovací. Akt o umělé inteligenci by proto neměl vést ke vzniku nadměrné zátěže pro startupy, které v oblasti AI působí. Startupy totiž nemusí mít dostatečnou finanční a administrativní kapacitu k tomu, aby dokázaly splnit přísné požadavky na testování.
- Pokud bude na základě článku 1 zavedeno preventivní testování pro vysoce rizikové modely AI (což podporujete), navrhuje výjimku pro startupy s méně než 20 zaměstnanci a obratem do 4 milionů eur. V jejich případě by členské státy raději měly uplatňovat flexibilní přístup. Považujete to za vyvážený kompromis a chcete o něm vést diskusi.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

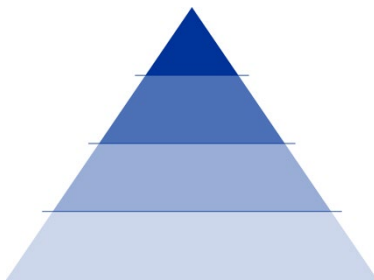
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabíjet bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie	na základě schopností	předběžná opatrnost		startupy	
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

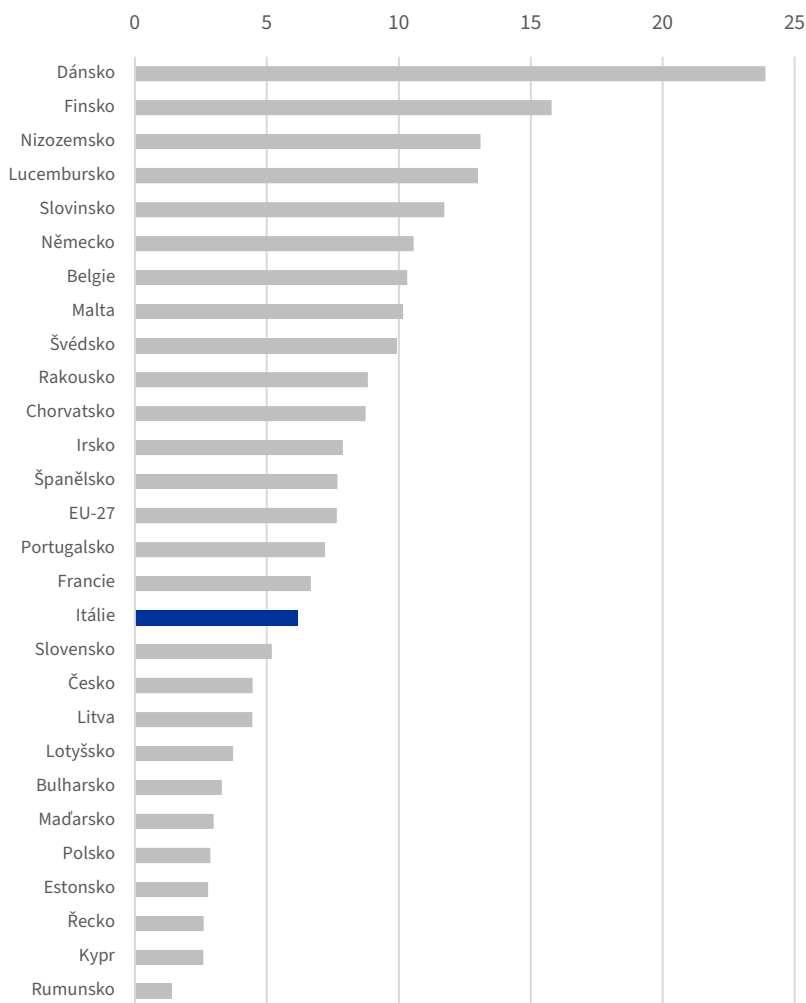
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Itálie
Počet obyvatel	447 695 350	58 997 201
HDP na obyvatele (v EUR)	38 150	36 130
Míra nezaměstnanosti	6,1 %	7,7 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	5,1 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	22,2 %

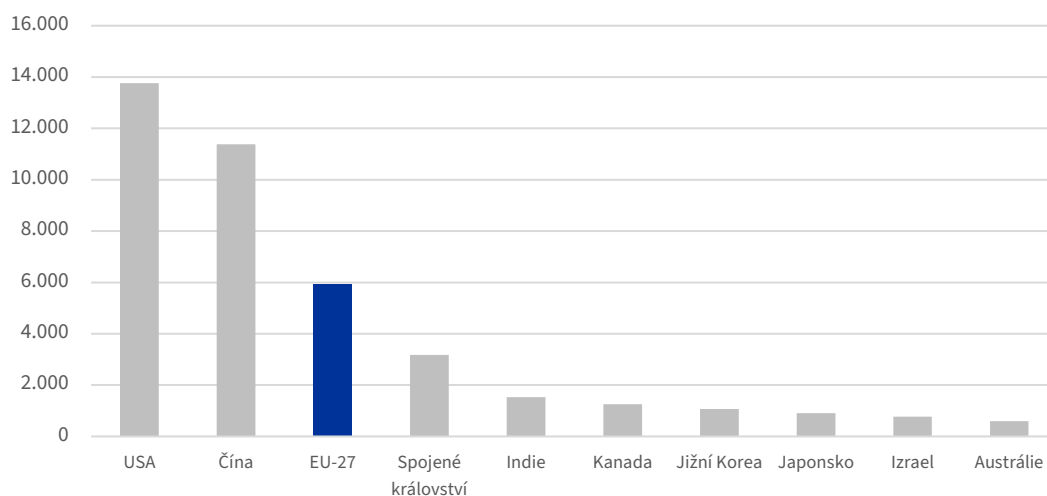
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



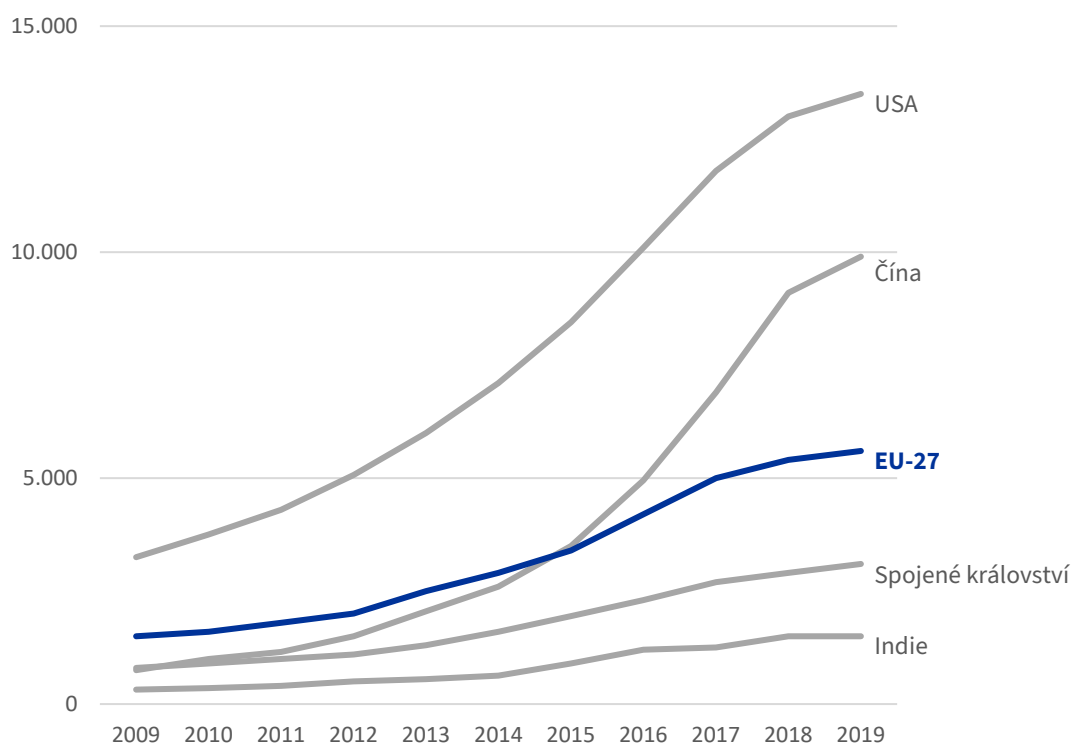
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



LITVA



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Česko, Estonsko, Lotyšsko a Maďarsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Výroba, jež je největším odvětvím litevské ekonomiky, generuje 20,4 % HDP této země. Jednou z největších výzev, kterým litevské odvětví výroby čelí, je nízká produktivita práce. Prostřednictvím automatizace rutinních činností a optimalizací procesů může s řešením tohoto problému pomoci AI.
- Zasazujete se proto o zavedení právního rámce, který umožní bezplatné inovace a experimentování. I když je třeba zabránit nechtěným důsledkům, zdůrazňujete, že v rámci flexibilního přístupu jsou i nadále povinná vlastní hodnocení. Vývojáři je však mohou provést individuálně způsobem, jenž bude pro jejich produkty nejvhodnější.
- Přílišná regulace by zpomalila inovace a poškodila evropské startupy v oblasti AI. Povinné externí audity a zdoluhavé testovací fáze, které právní předpisy vyžadují, totiž znamenají vysoké náklady.
- Ačkoli některé členské státy prosazují přísnější pravidla kvůli obavám ohledně lidských práv, z vašeho pohledu může být AI také účinným nástrojem, který bude tato práva chránit. AI může pomoci snižovat předpojatost při náboru pracovníků, policejní práci a rozhodování, ale také ji lze využít k obraně proti kybernetickým útokům nebo k ochraně osobních údajů. I když může AI vytvářet neúmyslně nesprávné informace nebo tzv. deepfakes, je možné ji také vytrénovat tak, aby tyto jevy odhalovala a bojovala proti nim. Inovace zaměřené na využití AI k ochranným účelům by měla EU aktivně podporovat.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Přestože prosazujete, aby byl článek 1 méně restriktivní a poskytoval větší svobodu pro inovace, je podle vás stále důležité mít stanovená jasná pravidla. Regulace má největší smysl, pokud platí pro všechny. Vývojářům poskytuje jasnější představu o tom, co se od nich očekává, a EU pomáhá utvářet jednotný přístup k odpovědné AI.
- V souvislosti s využitím AI už existuje řada varovných signálů. V Číně byla použita technologie AI na rozpoznávání obličeje a emocí s cílem identifikovat příslušníky muslimské menšiny Ujgurů, což vedlo k autocenzuře a k umístění řady osob do tzv. převýchovných táborů. V USA použila společnost Clearview AI bez předchozího souhlasu miliardy obrázků ze sociálních sítí a svůj nástroj prodala například agentuře FBI nebo Úřadu pro přistěhovalectví a výběr cel (ICE). Ve Spojeném království pak Metropolitní policie využila při své práci technologii rozpoznávání obličeje v reálném čase, nicméně výsledky vykazovaly vysokou míru chyb, zejména u osob s tmavou pletí.
- Podle vás by těmto problémům bylo možné zabránit řádným testováním a zavedením pravidel, která se budou vztahovat na všechny osoby bez jakýchkoli výjimek.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

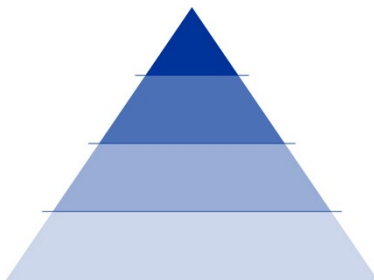
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

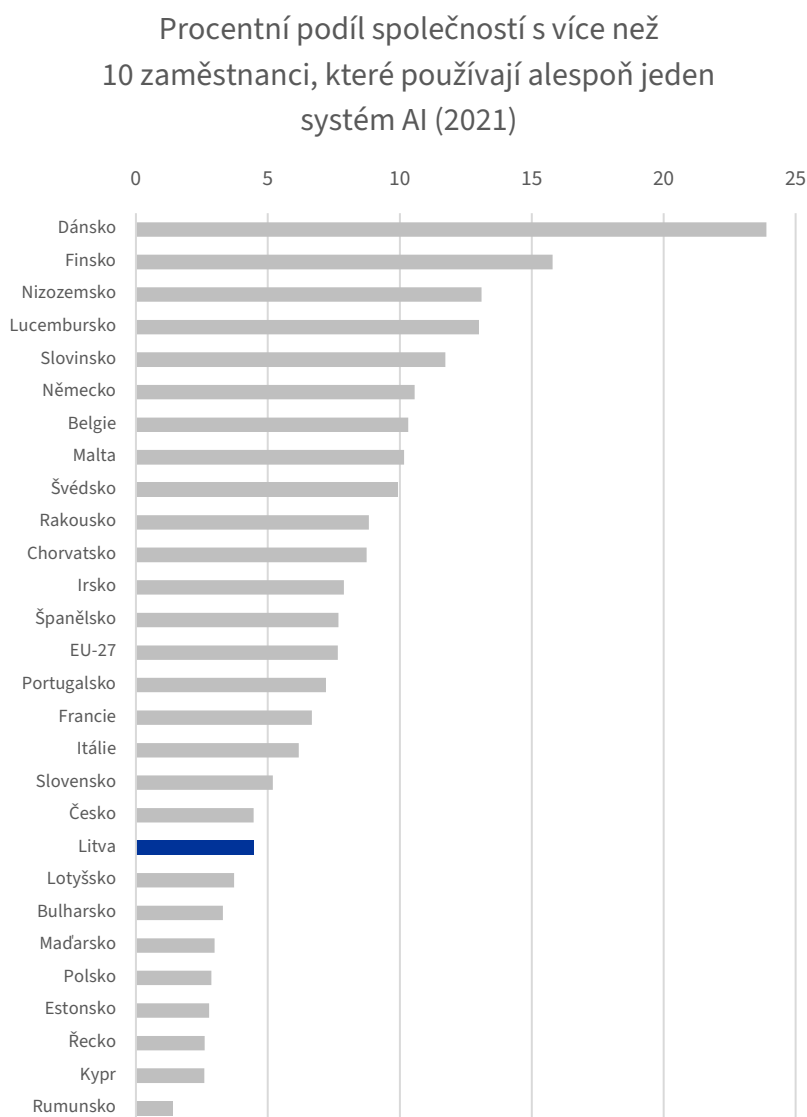
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva	na základě schopností	flexibilní		bez výjimky	
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

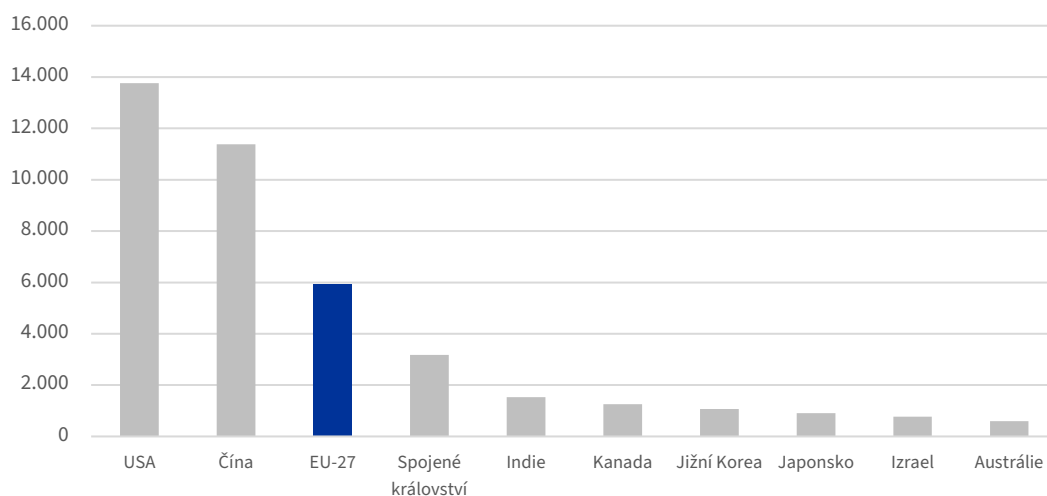
	Evropská unie	Litva
Počet obyvatel	447 695 350	2 857 279
HDP na obyvatele (v EUR)	38 150	25 700
Míra nezaměstnanosti	6,1 %	6,9 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	4,9 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	25,9 %



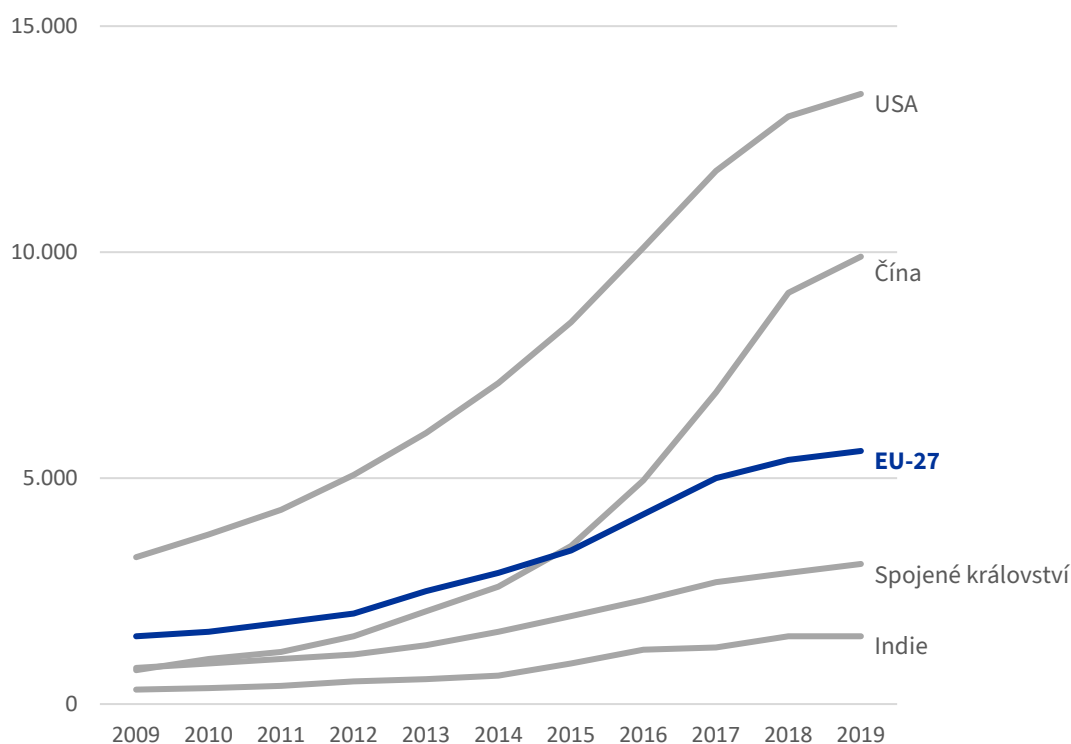
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



LUCEMBURSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **prečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Bulharsko, Malta a Slovensko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost do **tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Země, jako je Německo a Španělsko, se snaží z AI odstranit veškerou předpojatost, což považujete za nerealistické. Každé rozhodnutí, které děláme, s sebou nese určitou mírou předpojatosti. Pokud je AI vyvíjena odpovědně, může předpojatost u lidí omezovat. Snažíte se nasměrovat debatu tak, aby byla realističtější a odrážela potřeby všech členských států. Pokud vám někdo vytýká, že rizika spojená s AI zlehčujete, rozhodně nesouhlasíte – z vašeho pohledu je třeba klást důraz na rizika, která jsou skutečně bezprostřední. Z tohoto důvodu je podle vás nezbytný přístup založený na účelu.
- AI například nesmí představovat nebezpečí pro elektrické nebo internetové sítě, aby nedošlo k přerušení provozu nemocnic a nebyly ohroženy životy v řadě regionů. Přístup založený na schopnostech je příliš matoucí a dává prostor ke vzniku řady nejasností.
- Jste pro rychlejší rozvoj AI, a to i ve vysoce rizikových odvětvích. Obáváte se však, že vaše země nemá dostatek finančních prostředků k tomu, aby potenciál AI dokázala, vzhledem k přísným pravidlům, plně využít. Proto se zasazujete o zavedení určitých výjimek z požadavků, jejichž seznam je uveden v příloze II. Chcete, aby existovala možnost testování v reálných podmínkách za předpokladu, že externí audity zůstanou povinné. Tím byste umožnili, aby se vývojáři mohli na základě svých finančních možností rozhodnout, jaký druh testování zvolí.
- Pokud je obtížné dosáhnout shody na přístupu založeném na účelu, **jste ochotni souhlasit s doplněním ustanovení o přezkumu po třech letech.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- AI a digitální revoluce musí podle vás také podporovat genderovou rovnost v EU. V odvětvích souvisejících s výzkumem a inovacemi jsou ženy nedostatečně zastoupeny téměř ve všech oblastech a na všech úrovních. Zdůrazňujete, že je důležité aktivně podporovat zapojení žen do rozhodování o dalším formování AI.
- Startupy představují jedinečnou příležitost. Na rozdíl od velkých technologických společností existují ve startupech často menší genderové rozdíly. Mimo to i v oblasti AI nabízejí startupy více možností pro ženy a genderqueer osoby. Je také prokázáno, že různorodé týmy budují méně předpojatou a etičtější umělou inteligenci – něco, co by EU měla podporovat.
- Ze své podstaty usilují startupy o rychlé inovace s omezenými zdroji. Tento potenciál může potlačit velká administrativní zátěž, přičemž tomuto riziku jsou vystaveny zejména malé týmy, které se pokoušejí prosadit.
- Proto podporujete, aby malé startupy (s méně než 10 zaměstnanci a obratem do 2 milionů eur) byly vyňaty z plného rozsahu nařízení EU o umělé inteligenci.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

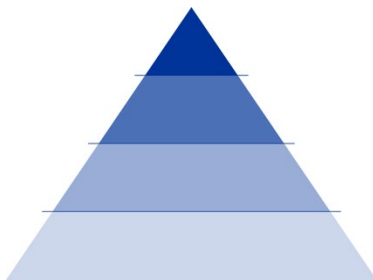
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko	na základě účelu	flexibilní	přezkum za 3 roky	startupy	
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

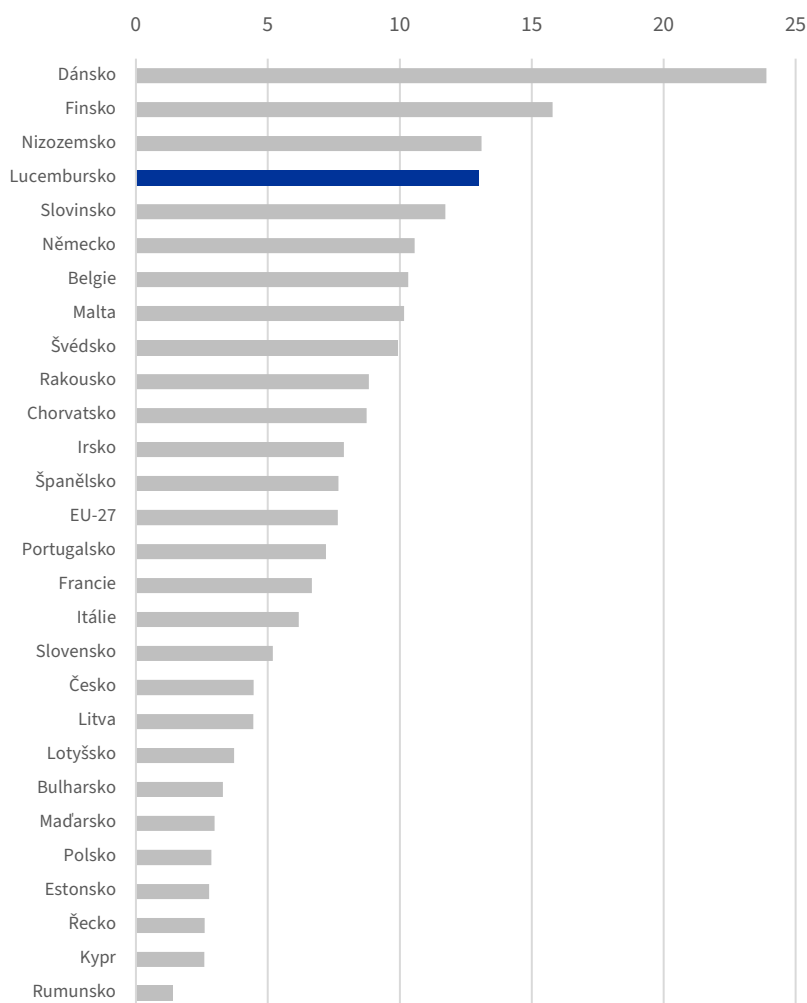
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Lucembursko
Počet obyvatel	447 695 350	660 809
HDP na obyvatele (v EUR)	38 150	121 290
Míra nezaměstnanosti	6,1 %	5,2 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	14,5 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	27,9 %

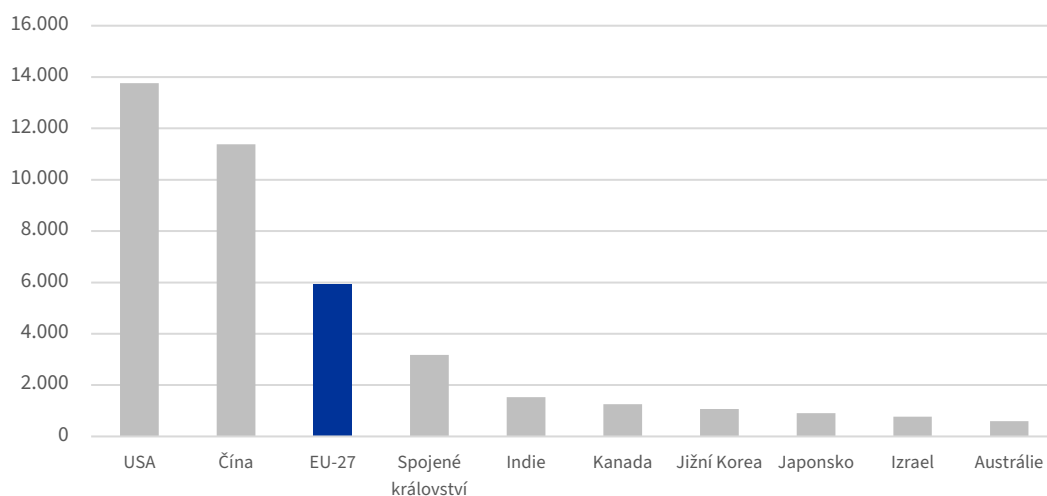
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



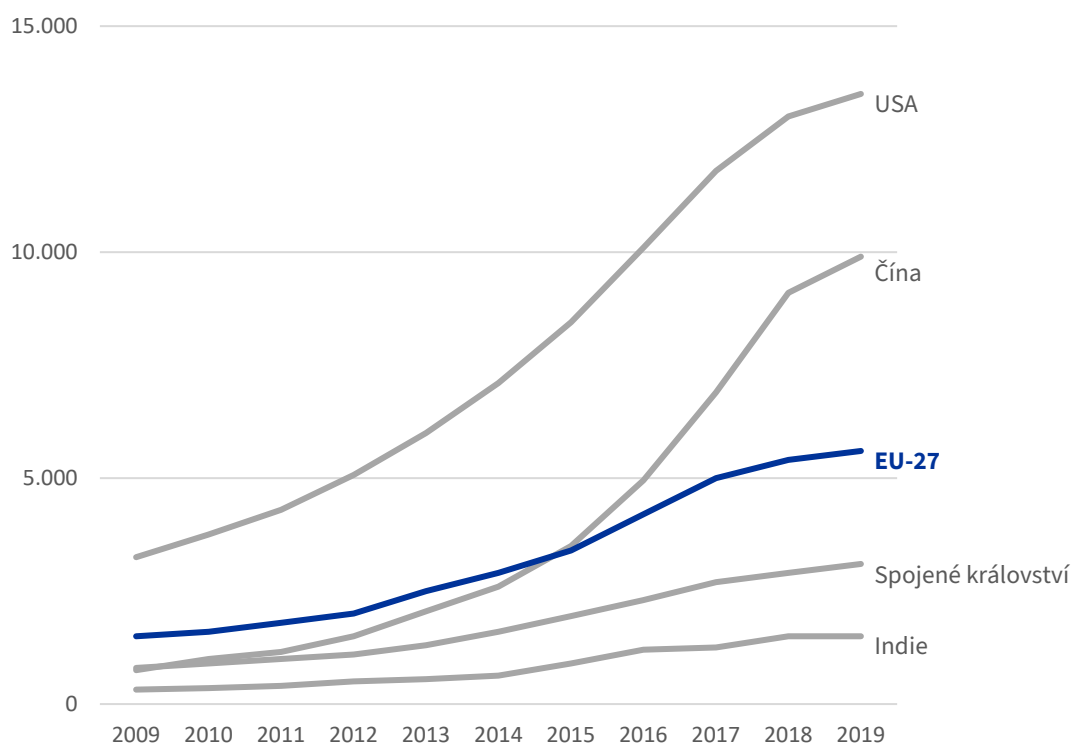
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



LOTYŠSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se zájmů a postojů vaší skupiny.**
2. **Některé z vašich argumentů s vámi sdílají země, jako je Česko, Estonsko, Litva a Maďarsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uveďte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je pro přehlednost do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Jste odhodláni zajistit to, aby byla v Lotyšsku zavedena digitální řešení. Tím vznikají nové požadavky na ochranu citlivých informací a digitální infrastruktury, které jsou nezbytné pro zachování životně důležitých společenských funkcí.
- Pokud jde o klasifikaci „vysoce rizikových“ modelů AI, považujete přístup založený na schopnostech vhodnější než přístup založený na účelu, který navrhuje Komise. Je-li stejný systém AI považován v jednom kontextu za nízkorizikový a v jiném za vysoce rizikový, může to pro uživatele, vývojáře a investory představovat nejasnosti a právní nejistotu.
- Pokud se mají do oblasti AI přilákat nezbytné investice, je nutné nastavit vhodná pravidla. To je podle vás způsob, jak důkladně chránit základní práva a zároveň umožnit svobodné a flexibilní inovační procesy. Pokud chce EU konkurovat velmocím v oblasti AI, jako jsou USA a Čína, musí prokázat svoji excelenci. Té však za příliš restriktivních pravidel, která omezují experimentování a vývoj, dosáhne jen stěží. Z těchto důvodů byste chtěli navrhnout určité výjimky ze seznamu požadavků na vysoce rizikové systémy AI (viz příloha II). Podle vás **by bylo nejlepší z povinných požadavků odstranit „monitorování po uvedení na trh“**, není totiž jasné, co přesně tato povinnost znamená.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Obecně souhlasíte se zeměmi, které nechtějí žádné výjimky, včetně vašeho souseda a silného spojence – Litvy. Je podle vás rozumné podporovat společný evropský rámec pro regulaci AI bez zbytečných výjimek.
- Od ruské invaze na Ukrajinu jste však spolu s ostatními pobaltskými státy ve stavu pohotovosti z důvodu možného ruského útoku na Evropskou unii. Dokonce máte s Ruskem společnou hranici, kterou každodenně chráníte spolu se silami NATO.
- Proto jste ochotni jednat o omezených výjimkách pro vojenské a obranné účely, nikoli však pro národní bezpečnost v širším slova smyslu. Přestože národní bezpečnost patří do pravomoci členských států, ponechat ji zcela mimo rámec EU by mohlo vést ke vzniku právních mezer a roztříštěnosti právních předpisů. Alespoň částečné zapojení EU může podpořit soudržnost (a srozumitelnost pro vývojáře) a dát najevo jednotný unijní postoj k AI v oblasti národní bezpečnosti.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

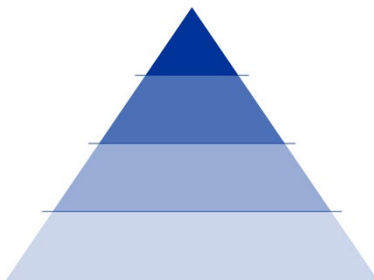
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko	na základě schopností	flexibilní	odstranit požadavek na monitorování po uvedení na	armáda/obrana	
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

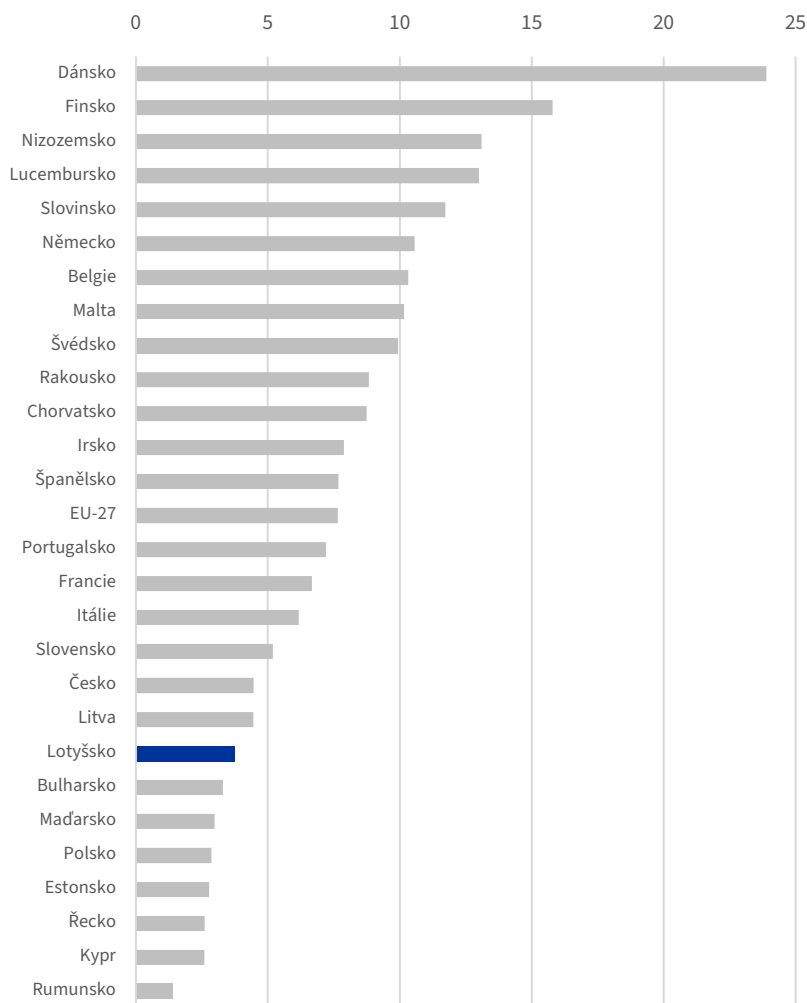
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Lotyšsko
Počet obyvatel	447 695 350	1 883 008
HDP na obyvatele (v EUR)	38 150	20 930
Míra nezaměstnanosti	6,1 %	6,5 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	4,5 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	16,6 %

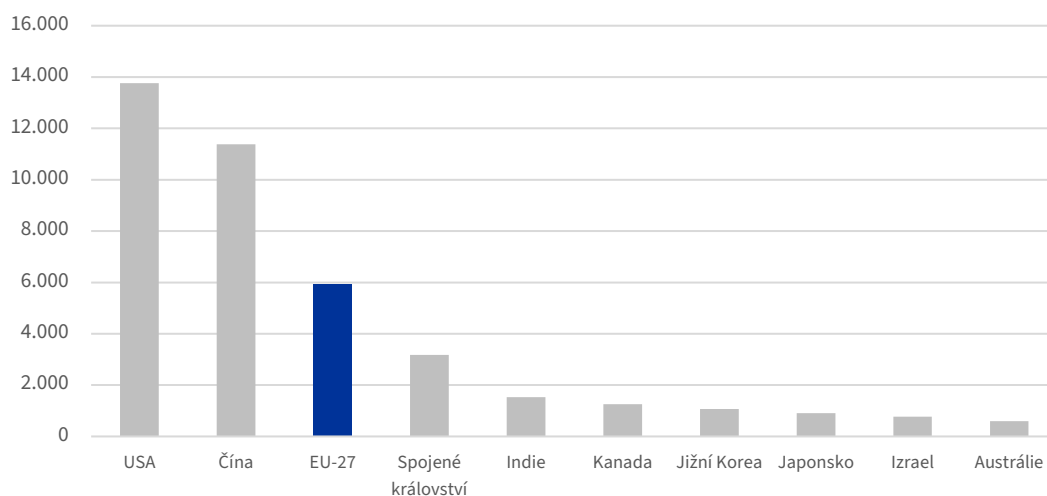
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



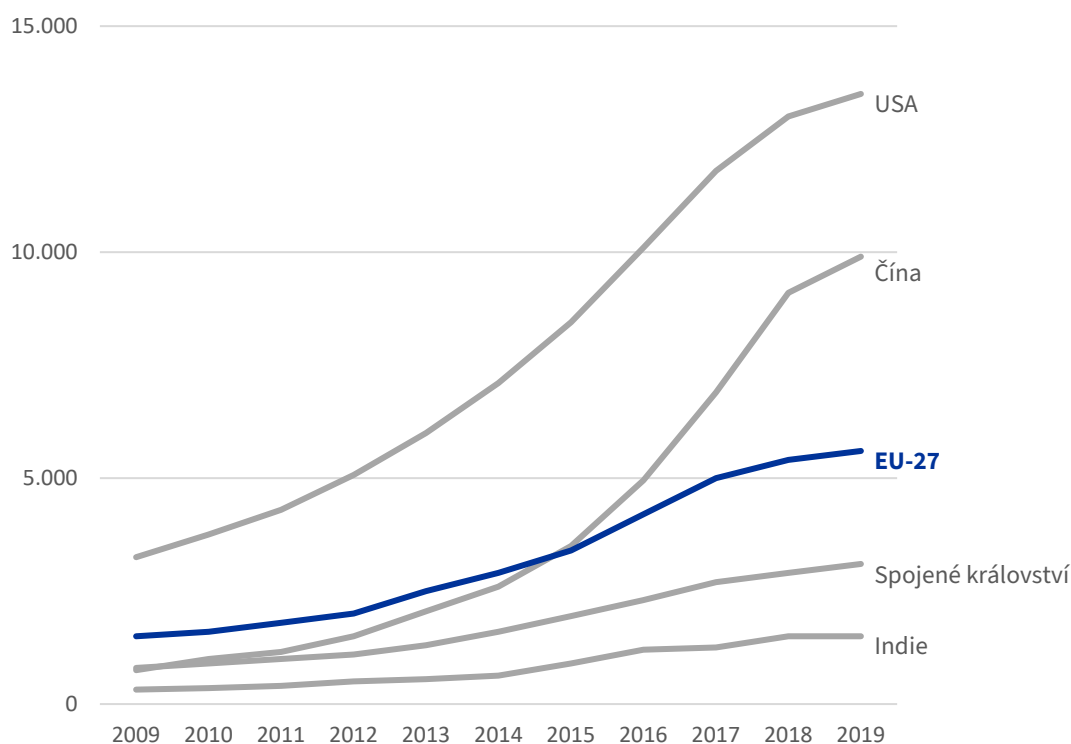
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



MALTA



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



MALTA

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **prečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Bulharsko, Lucembursko a Slovensko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Jako jedna z nejmenších členských zemí EU s pouze 550 000 obyvatel a jako ostrovní stát jste pevně odhodláni podporovat udržitelný hospodářský růst. Snažíte se přilákat zahraniční investice a kvalifikované talentované pracovníky, abyste zajistili rozvoj silného odvětví AI se zaměřením na medicínu a kritickou infrastrukturu. To jsou oblasti, u nichž se očekává, že přinesou vysoké výnosy.
- Podporujete návrh Komise týkající se horizontálního rámce pro AI, a zejména její přístup k řízení rizik založený na účelu, který považujete za logický a vhodný, protože pro regulační orgány je snazší posuzovat konkrétní aplikace než hodnotit abstraktní schopnosti modelu.
- V současné době je nutné zdůraznit, že technologická kreativita vyžaduje svobodu – inovace je třeba podporovat, nikoli omezovat. Z toho důvodu **je důležité prosazovat testování nových modelů AI v reálných podmínkách**, neboť prostředí testovacího pískoviště, i když je užitečné, často reálným podmínkám plně neodpovídá.
- Malé země, jako je Malta, čelí specifickým výzvám, jako je omezený počet talentovaných pracovníků a méně zdrojů na nákladné testování. Naopak větší státy s velkým kapitálem se mohou s touto zátěží vypořádat lépe. Důrazně proto nesouhlasíte s příliš přísnými požadavky na testování, které by systematicky bránily dosažení cílů vaší země v oblasti vývoje AI.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- I nadále je pro vás zásadní, aby se nařízení o AI nevztahovalo na záležitosti týkající se národní bezpečnosti, armády nebo obrany. Zejména národní bezpečnost není v pravomoci EU. V ustanovení čl. 4 odst. 2 Smlouvy o Evropské unii (SEU) se uvádí: „...národní bezpečnost zůstává výhradní odpovědností každého členského státu.“ Proto nemá smysl členské státy při regulaci AI v aplikacích určených k zajištění národní bezpečnosti omezovat.
- Na druhé straně nevidíte pádný důvod pro to, aby se nařízení nevztahovalo na startupy. To, že je společnost malá, ještě neznamená, že její technologie nemohou způsobit škodu. Pokud by se navíc na startupy vztahovala výjimka, mohlo by se stát, že velké společnosti budou startupům zadávat rizikové projekty, aby se vyhnuly požadavkům na testování. Vznikla by tak mezera, které by mohl někdo využít, a EU by měla zajistit, aby k tomu nedošlo.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

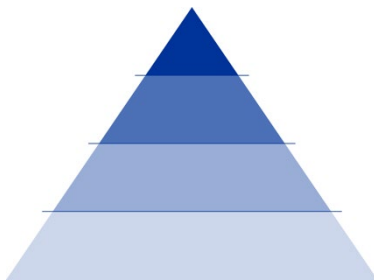
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta	na základě účelu	flexibilní	testování v reálných podmínkách	národní bezpečnost	
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

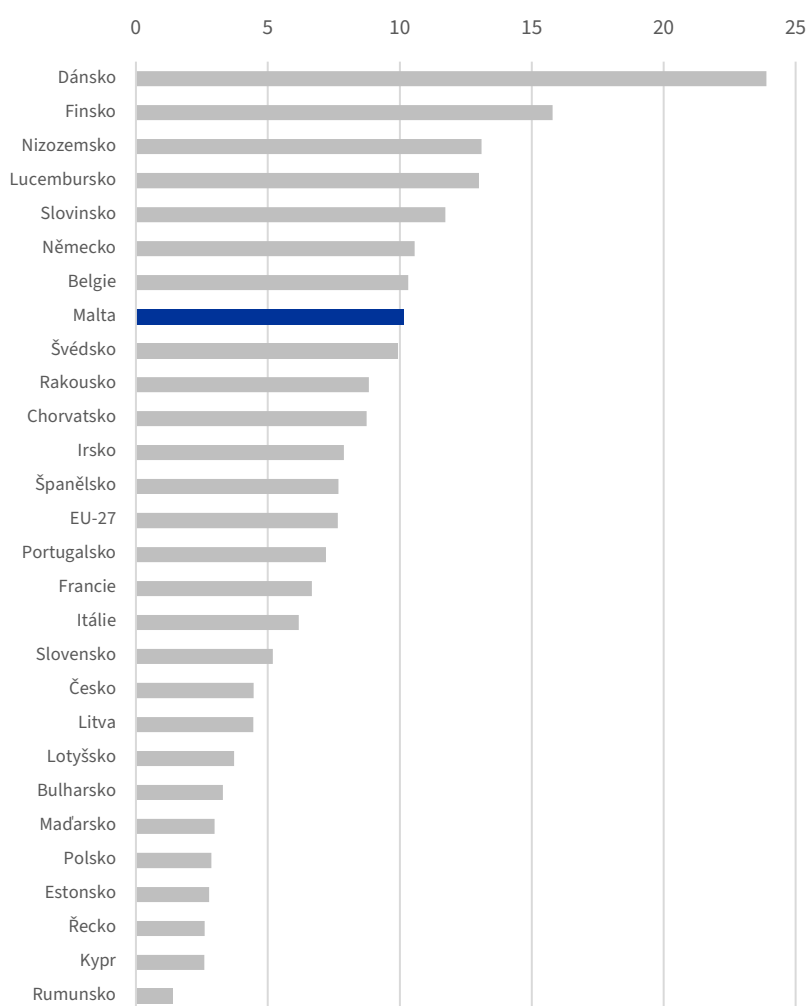
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Malta
Počet obyvatel	447 695 350	542 051
HDP na obyvatele (v EUR)	38 150	37 110
Míra nezaměstnanosti	6,1 %	3,5 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	13,2 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	37 %

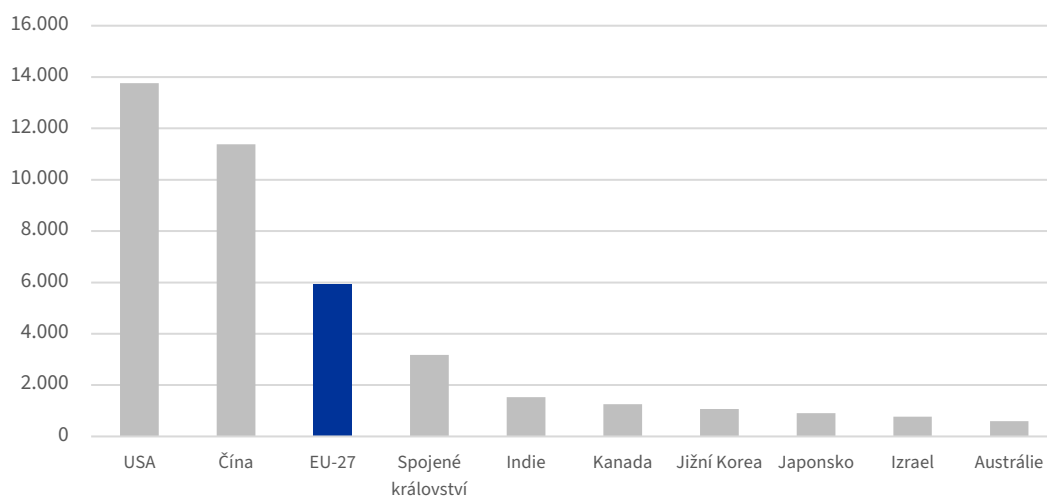
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



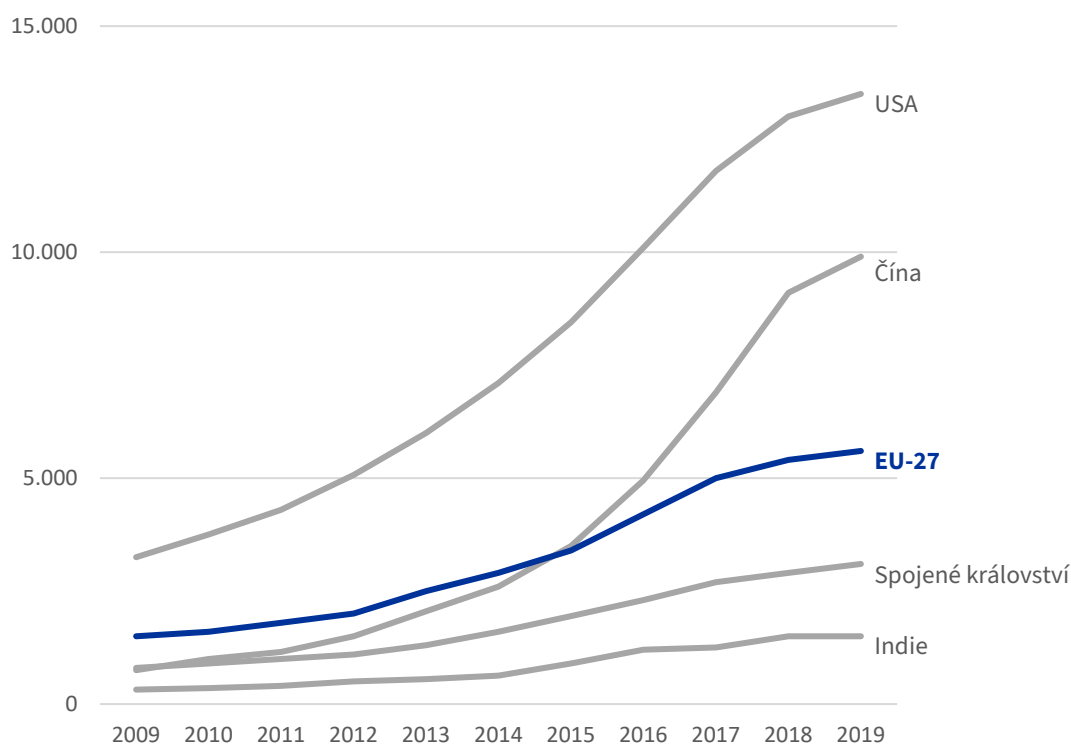
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



NIZOZEMSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **prečtete informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílají země, jako je Chorvatsko, Itálie a Slovinsko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.



*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Zachování národních hodnot a hospodářské prosperity je pro vládu vaší země nejvyšší prioritou. Podle Mezinárodního měnového fondu může mít AI dopad až na 60 % pracovních míst, v důsledku čehož se sníží jejich počet, a naopak se zvýší potřeba odborníků v oblasti AI.
- Nekontrolovaná nezaměstnanost by nizozemské hospodářství vážně ohrozila, a proto důrazně podporujete preventivní testování, aby se zabránilo selhání vysoce rizikových technologií a narušení celých odvětví.
- Jste také zastáncem zvyšování gramotnosti v oblasti AI napříč nizozemskou společností, aby byli lidé schopni rozpoznat předpojatost nebo manipulaci a slepě nedůvěřovali automatizovaným systémům. Přístup založený na schopnostech je pro tento cíl vhodnější, neboť se zaměřuje na skutečná rizika zneužití a manipulace. Některé schopnosti, jako je rozpoznávání obličejů, jsou ze své podstaty spojeny s vysokým rizikem zneužití nebo selhání, a to v jakékoli situaci.
- Systém založený na schopnostech zabraňuje nedostatečné regulaci těchto technologií jednoduše proto, že jsou používány v „nízkorizikových“ odvětvích. Například systém AI využívaný v odvětví maloobchodu, který sleduje pohyb zákazníků na prodejní ploše, předvídá jejich chování a používá rozpoznávání obličejů k personalizované nabídce slevových akcí nebo k označování „důležitých“ zákazníků, by se v rámci modelu založeného na odvětvích mohl kontrole vyhnout. Tyto základní schopnosti však vzbuzují vážné obavy ohledně soukromí, profilování a diskriminace.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Podporujete rámec preventivního testování podle článku 1, ale jste pro výjimky pro účely národní bezpečnosti. Národní bezpečnost není v pravomoci EU. V ustanovení čl. 4 odst. 2 Smlouvy o Evropské unii (SEU) se uvádí: „...národní bezpečnost zůstává výhradní odpovědností každého členského státu.“ Proto nemá smysl členské státy při regulaci AI v aplikacích určených k zajištění národní bezpečnosti omezovat.
- Na druhé straně nevidíte pádný důvod pro to, aby se nařízení nevztahovalo na startupy. To, že je společnost malá, ještě neznamená, že její technologie nemohou způsobit škodu. Pokud by se navíc na startupy vztahovala výjimka, mohlo by se stát, že velké společnosti budou startupům zadávat rizikové projekty, aby se vyhnuly požadavkům na testování. Vznikla by tak mezera, které by mohl někdo využít, a EU by měla zajistit, aby k tomu nedošlo.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

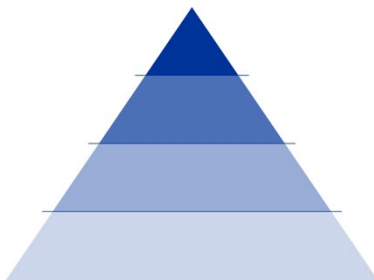
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabíjet bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko	na základě schopností	předběžná opatrnost		národní bezpečnost	
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

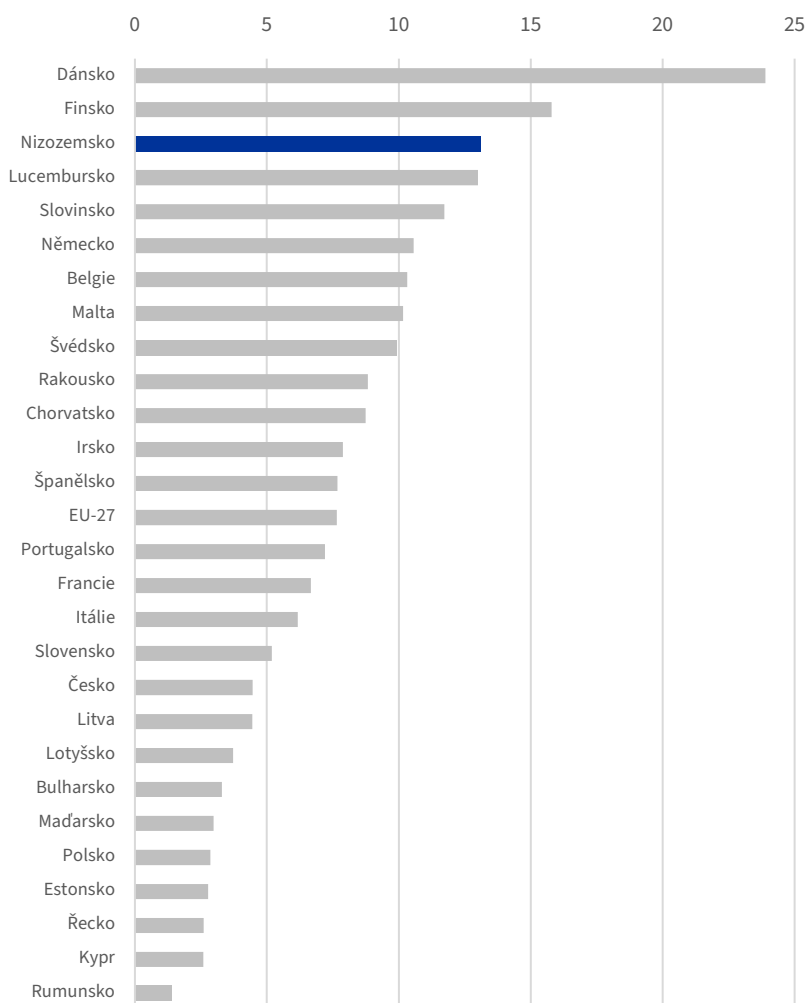
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Nizozemsko
Počet obyvatel	447 695 350	17 811 291
HDP na obyvatele (v EUR)	38 150	59 720
Míra nezaměstnanosti	6,1 %	3,6 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	14,1 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	54,5 %

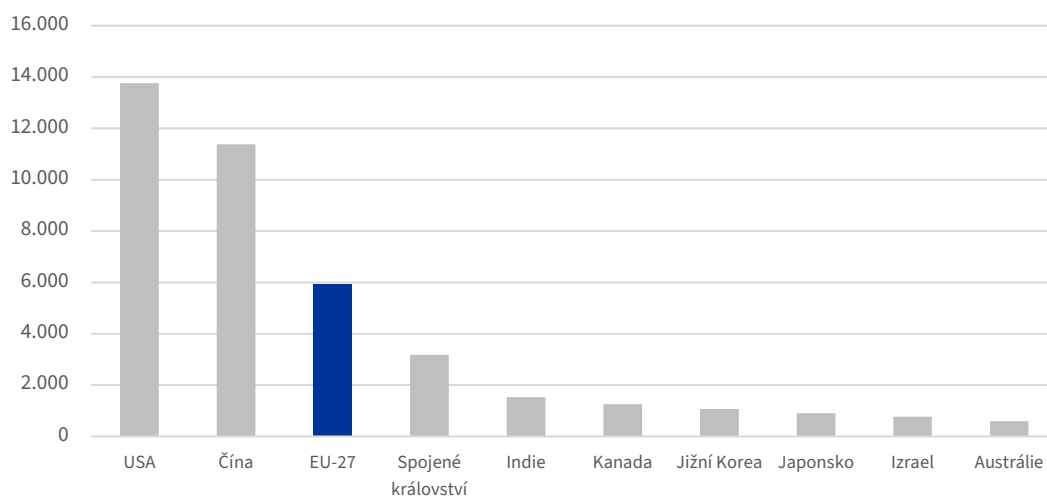
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



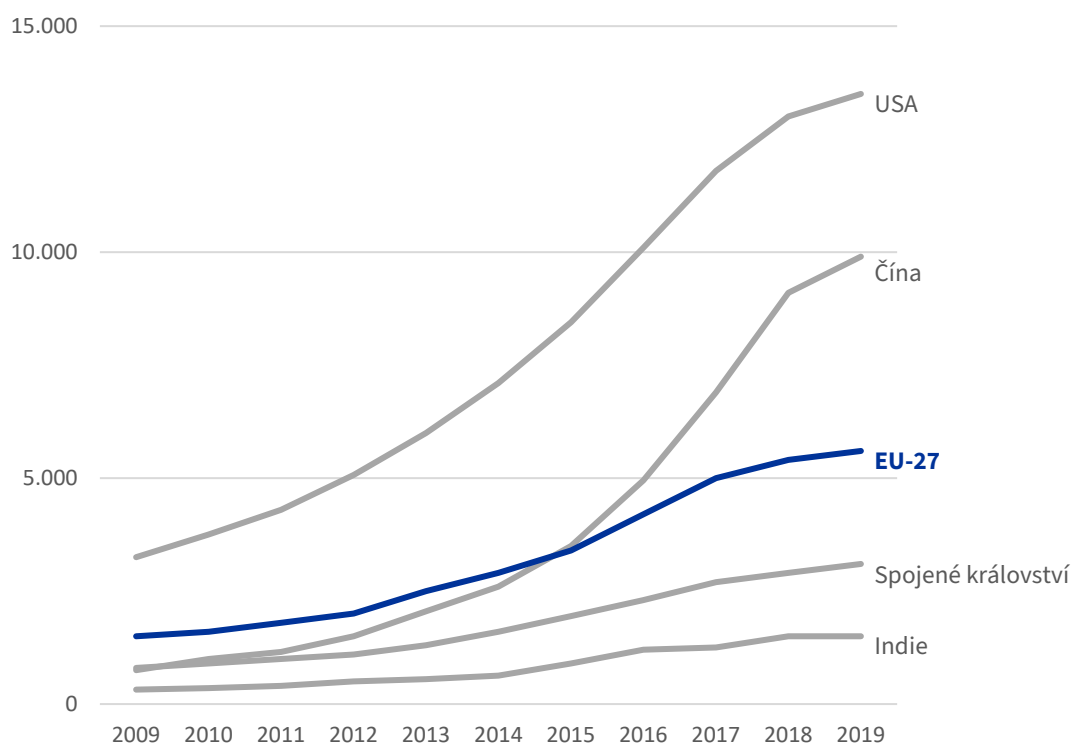
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

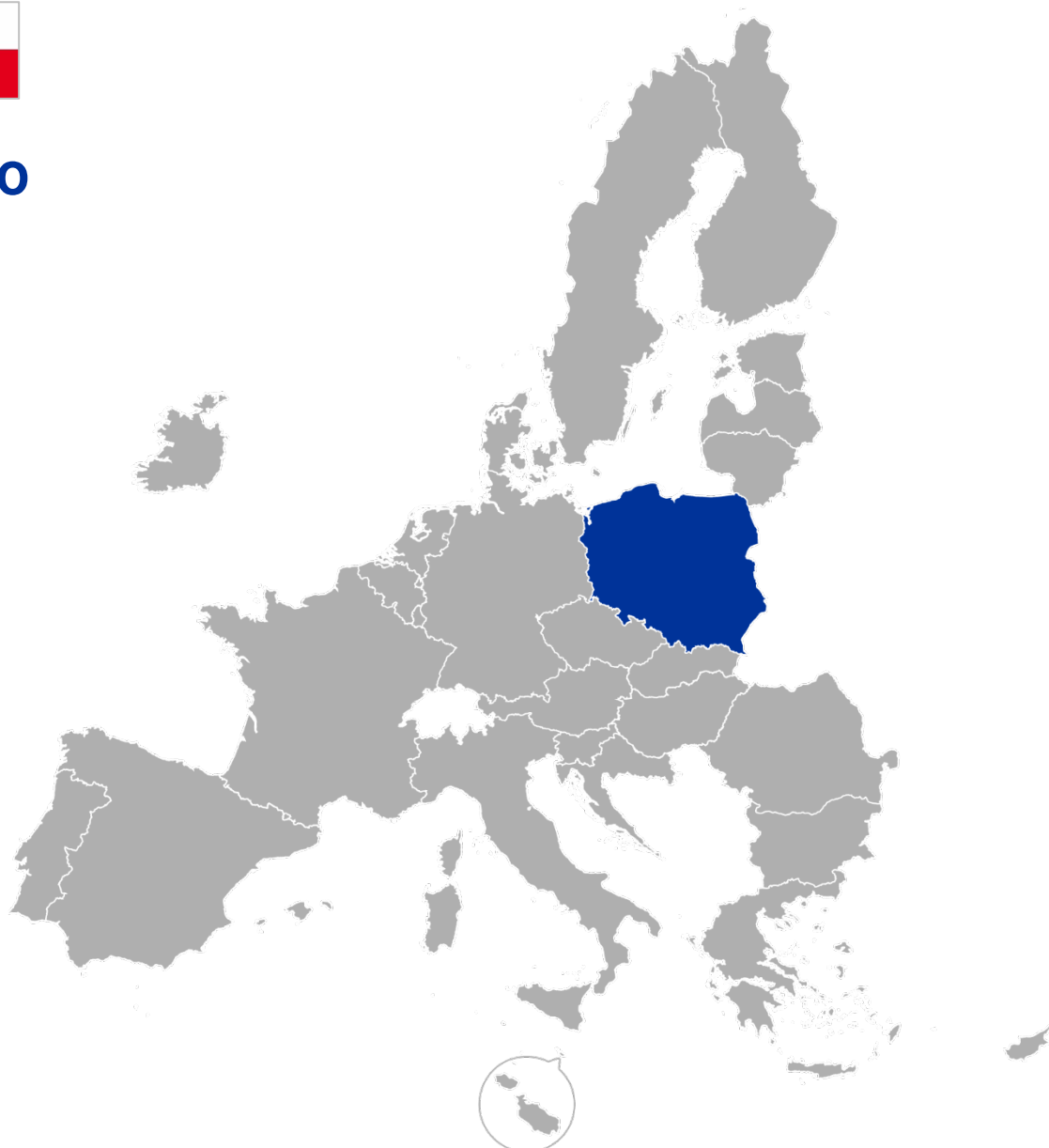
Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



POLSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



POLSKO

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Dánsko, Francie, Rakousko a Rumunsko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost do **tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Polsko má v současné době jen málo globálních technologických gigantů, proto je AI opravdovou příležitostí pro růst malých a středních podniků. Vaše ekonomika, která je založena na energetice, průmyslu, zemědělství a logistice, má dobré předpoklady k tomu, aby z používání AI těžila.
- Podporujete přístup založený na účelu, který vytváří jasné stimuly: nízkoriziková odvětví získají flexibilitu v oblasti inovací, zatímco vysoce riziková odvětví budou muset dodržovat přísné normy. Posílí se tak důvěra v evropskou AI, která je celosvětově známá pro svou kvalitu a bezpečnost.
- Uznáváte, že AI může představovat riziko předpojatosti, diskriminace a riziko v oblasti soukromí, protože žádný datový soubor není úplně neutrální. Přístup předběžné opatrnosti při testování zajišťuje, že se v případech, kdy je riziko újmy nejpravděpodobnější, použijí přísnější pravidla, což může pomoci zabránit jakémukoli negativnímu vlivu a chránit základní práva.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Prosazujete, aby se nařízení o umělé inteligenci nevztahovalo na aplikace ve vojenské a obranné oblasti. Zdůrazňujete, že nadměrná regulace v této oblasti by mohla bránit rychlému zavádění technologií AI v zónách konfliktu a během reakce na krizi. V souvislosti s přetrvávající nestabilitou v regionu má při používání těchto technologií zásadní význam flexibilita.
- Jako zemi, která sdílí hranici s Ruskem, Ukrajinou a Běloruskem, Polsko hluboce znepokojuje národní bezpečnost a bezpečnost jeho občanů. Podél hranice s Běloruskem navíc Polsko čelí značným výzvám, pokud jde o řízení zvýšeného pohybu uprchlíků.
- S ohledem na tyto tlaky usiluje Polsko o větší pochopení a podporu svých specifických geopolitických i bezpečnostních výzev ze strany ostatních členských států EU.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

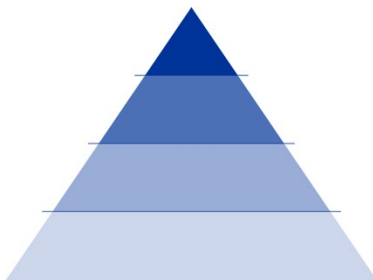
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

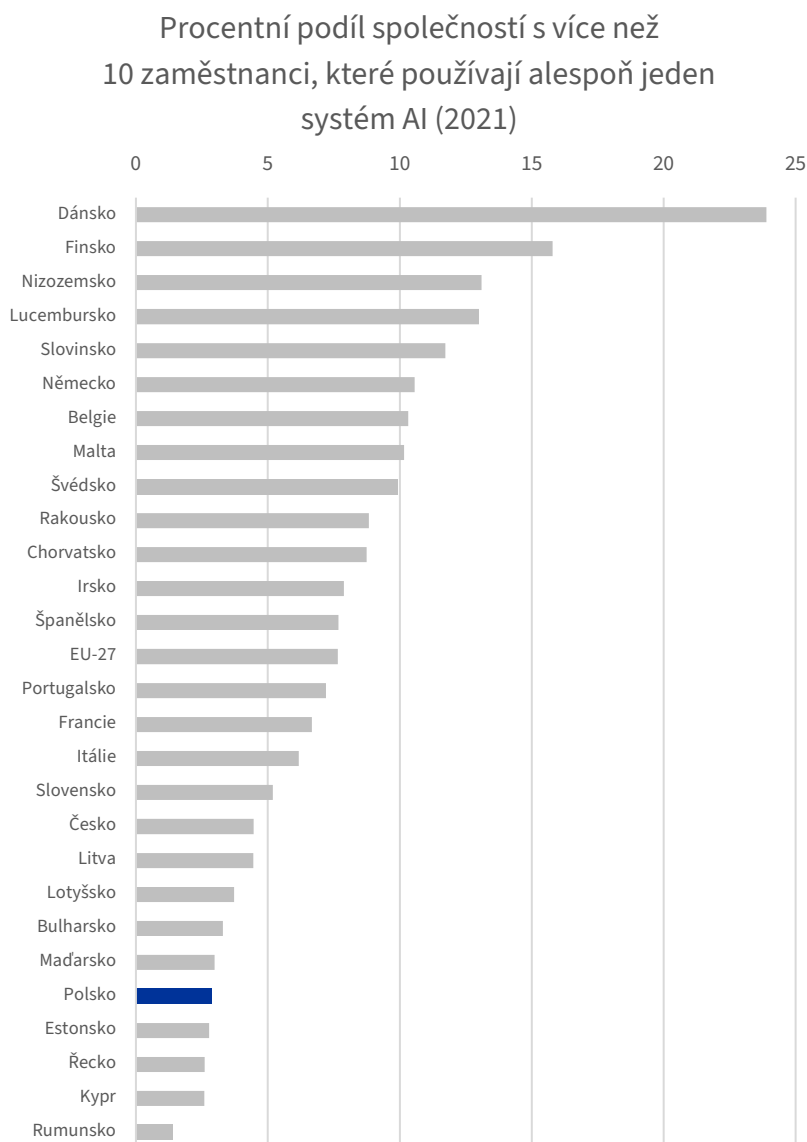
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko	na základě účelu	předběžná opatrnost		armáda/obrana	
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

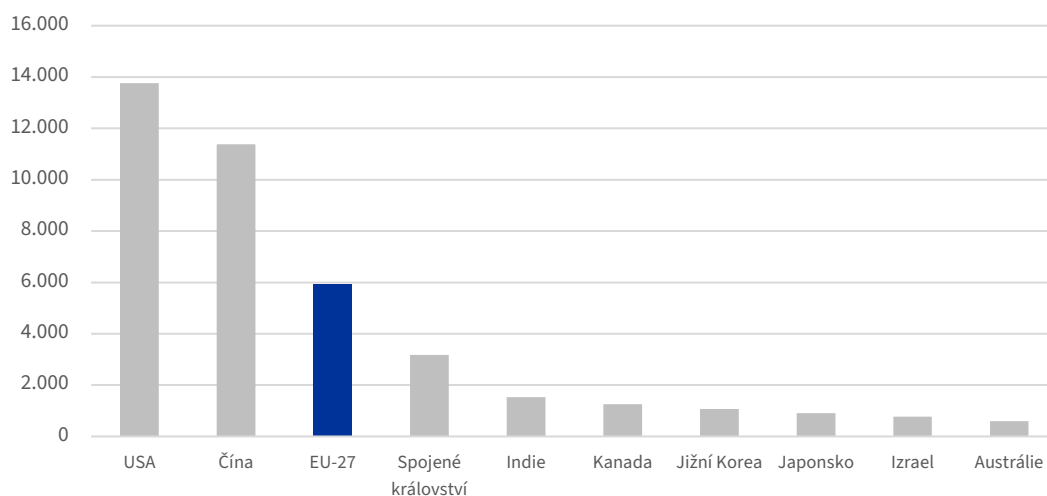
	Evropská unie	Polsko
Počet obyvatel	447 695 350	36 753 736
HDP na obyvatele (v EUR)	38 150	19 980
Míra nezaměstnanosti	6,1 %	2,8 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	3,7 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	20,1 %



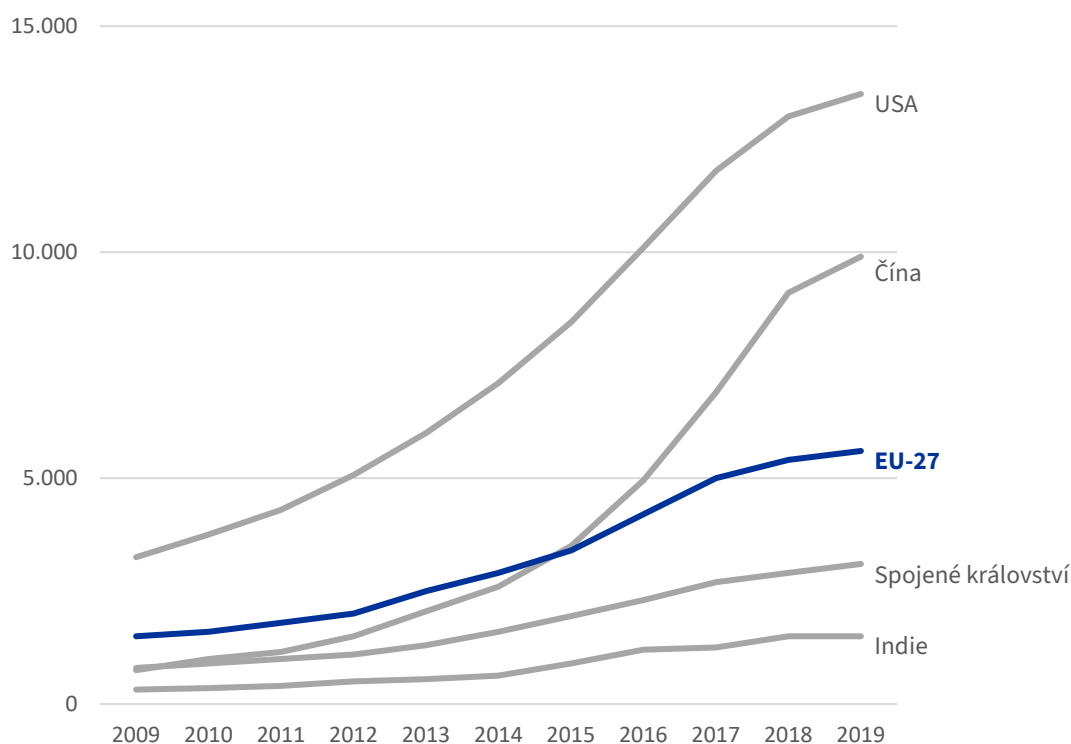
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

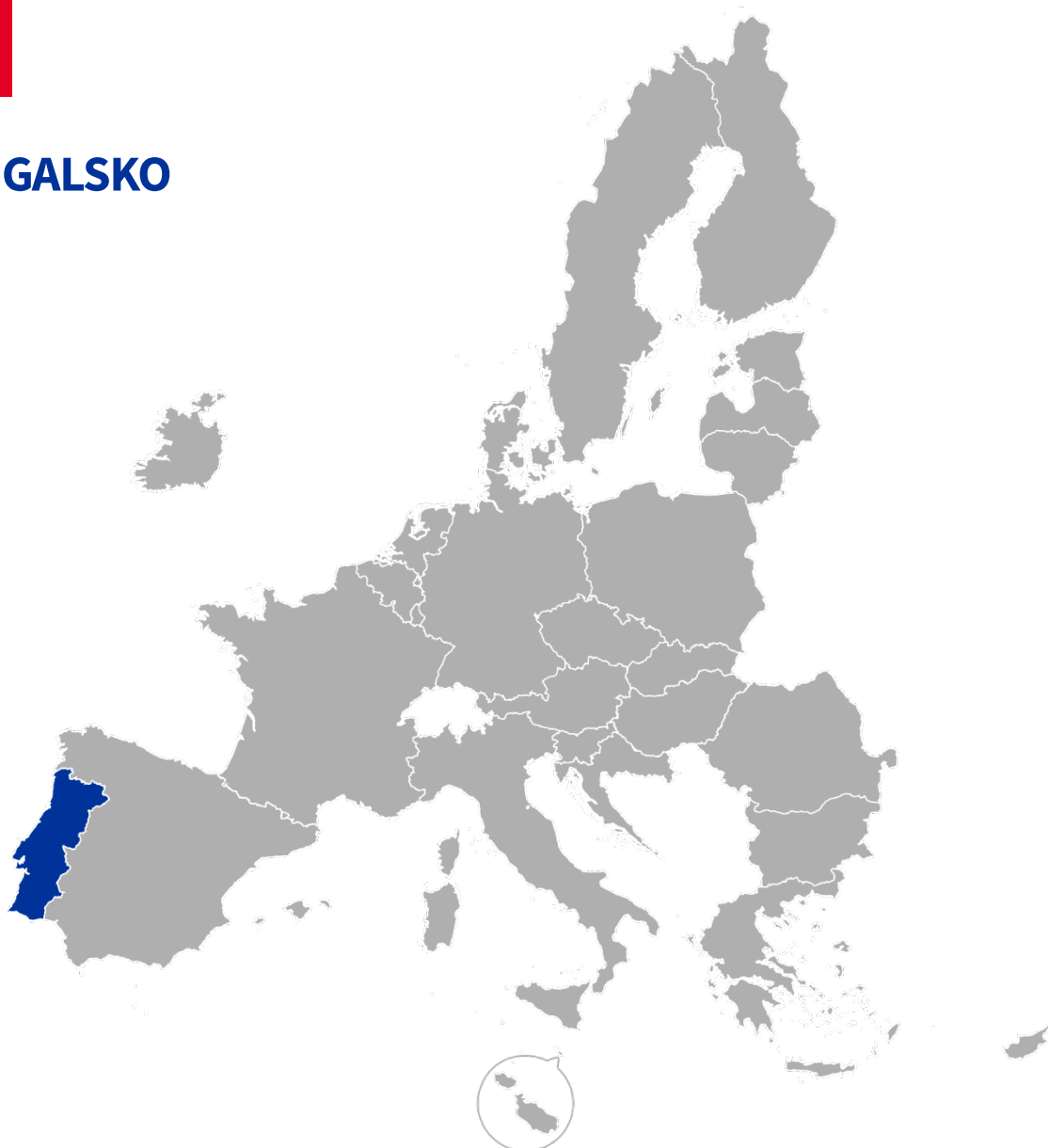
Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



PORTUGALSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílí země, jako je Finsko, Irsko a Švédsko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Portugalsko je vzhledem k technickým dovednostem svých pracovníků velmi atraktivní destinací pro mezinárodní high-tech společnosti. Jedním z příkladů je společnost Tekever, která se zaměřuje na využití dronů k autonomnímu sledování za pomoci AI pro obchodní i vojenské účely. Pokud jde o obchodní využití, pomocí těchto dronů lze sledovat potrubí, zda nedochází k únikům či sabotážím, a chránit tak kritickou infrastrukturu. Systémy Tekever využívá také Evropská agentura pro námořní bezpečnost (EMSA) při operacích pobřežní stráže a v rámci ochrany životního prostředí.
- Jste zastáncem přístupu založeného na účelu, ale obáváte se, že nepřesné popisy jednotlivých odvětví by mohly způsobit, že by akt EU o umělé inteligenci byl vnímán jako nedostatečný a nejasný. Chtěli byste zajistit, aby byl popis „kritické infrastruktury“ přesnější a například vyjasnit, že se jedná o výrobu elektřiny, telekomunikace, dodávky vody, zemědělství, vytápění a dopravu.
- Obáváte se, že preventivní testování by zvýšilo finanční a administrativní zátěž vývojářů AI.
- Doufáte, že se dnes dohodnete na konečném znění tohoto aktu, který bude vodítkem pro důvěryhodný vývoj AI v Evropě, zásadně ale nesouhlasíte s tím, aby byl v aktu stanoven přezkum po několika letech. Takové ustanovení by pouze narušilo důvěru investorů. Pokud by přetrvávala nejistota, **bylo by zodpovědnější odložit hlasování než preventivně počítat se změnou pravidel.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- I když podporujete flexibilní právní úpravu uvedenou v článku 1, jste i nadále přesvědčeni, že by měla platit jasná pravidla pro všechna odvětví – včetně národní bezpečnosti a zpravodajských služeb. Udělením výjimek těmto odvětvím vznikají právní mezery a zvyšuje se riziko toho, že nebude odhalena neetická nebo škodlivá AI.
- Cílem není zamezit využívání AI v oblasti národní bezpečnosti, ale zajistit, aby byla AI navržena zodpovědně tak, aby se předešlo předpojatosti nebo diskriminaci.
- Jedinou výjimkou, s níž můžete souhlasit, jsou modely s otevřeným zdrojovým kódem. Tyto modely významně urychlují inovace, protože jejich vývoj je jasnější. Vývojáři společně experimentují a vzájemně si vyměňují poznatky, díky čemuž je možné zdokonalovat nástroje. U AI s otevřeným zdrojovým kódem je často snazší provádět přezkumy a kontroly než u modelů s uzavřeným zdrojovým kódem, což je v souladu s cíli aktu.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

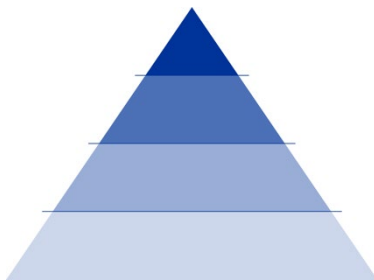
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

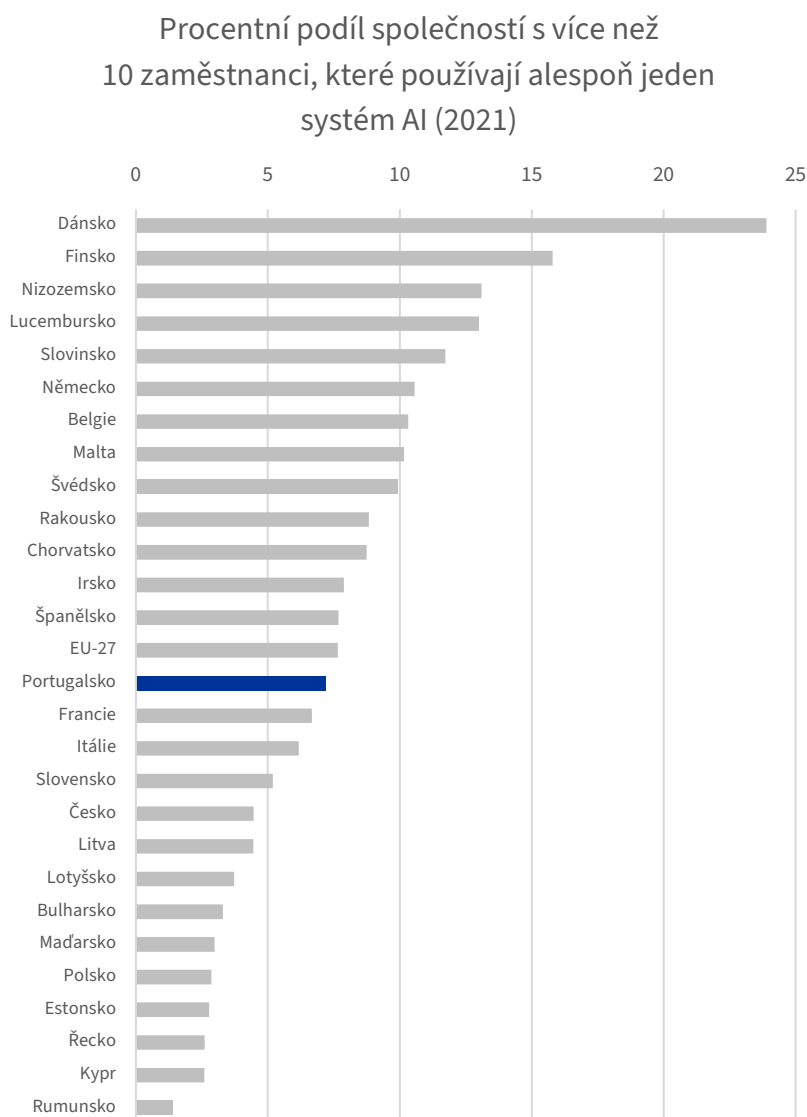
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko	na základě účelu	flexibilní	odložit hlasování	otevřený zdroj	
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

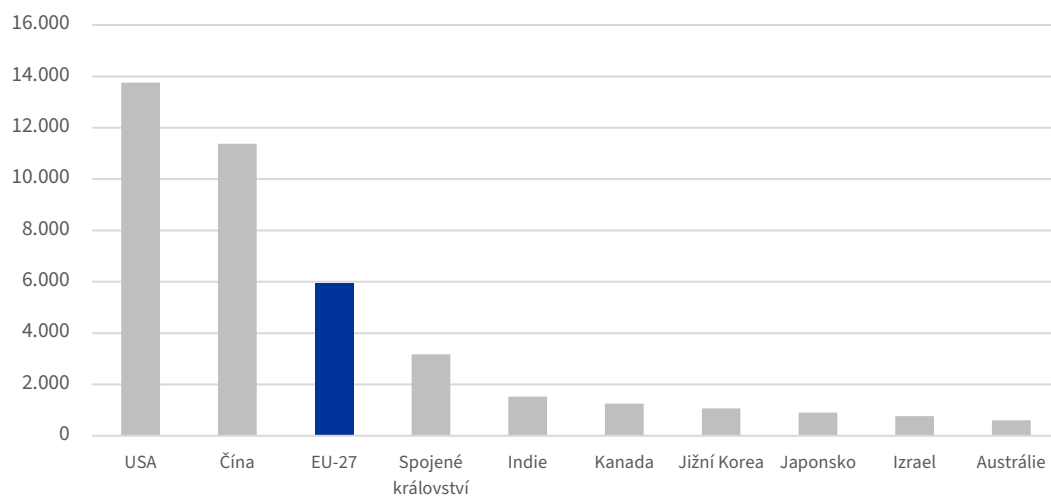
	Evropská unie	Portugalsko
Počet obyvatel	447 695 350	10 516 621
HDP na obyvatele (v EUR)	38 150	25 330
Míra nezaměstnanosti	6,1 %	6,5
Procentní podíl společností využívajících technologii umělé inteligence	8 %	7,9
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	29,9



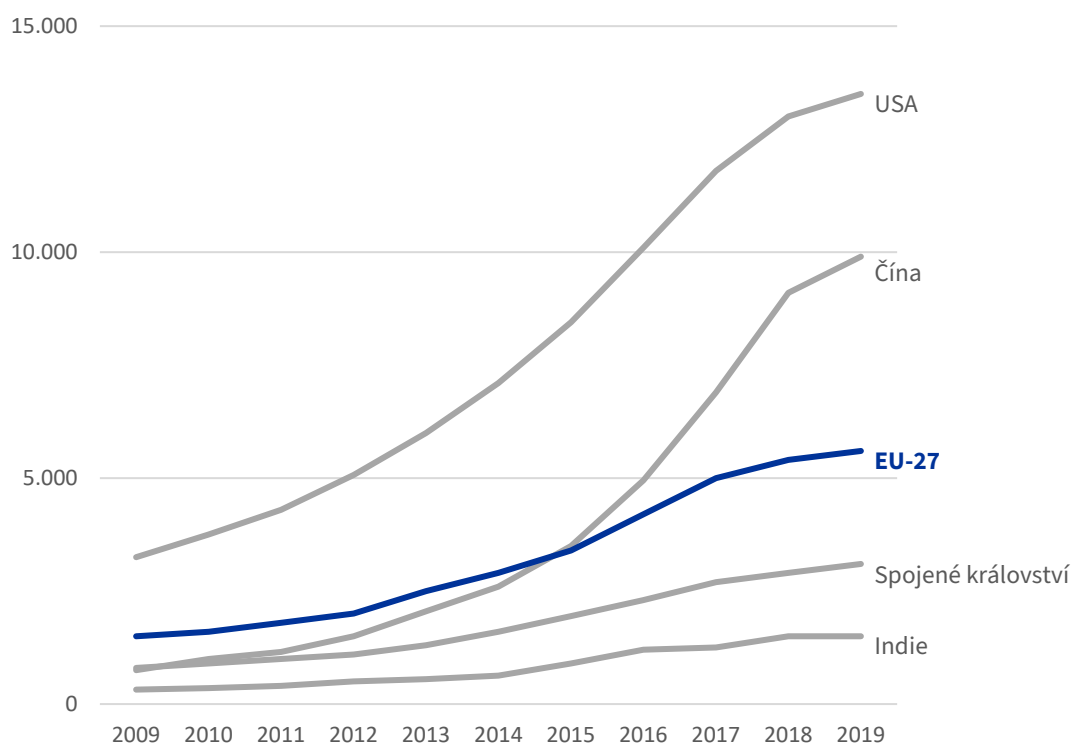
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

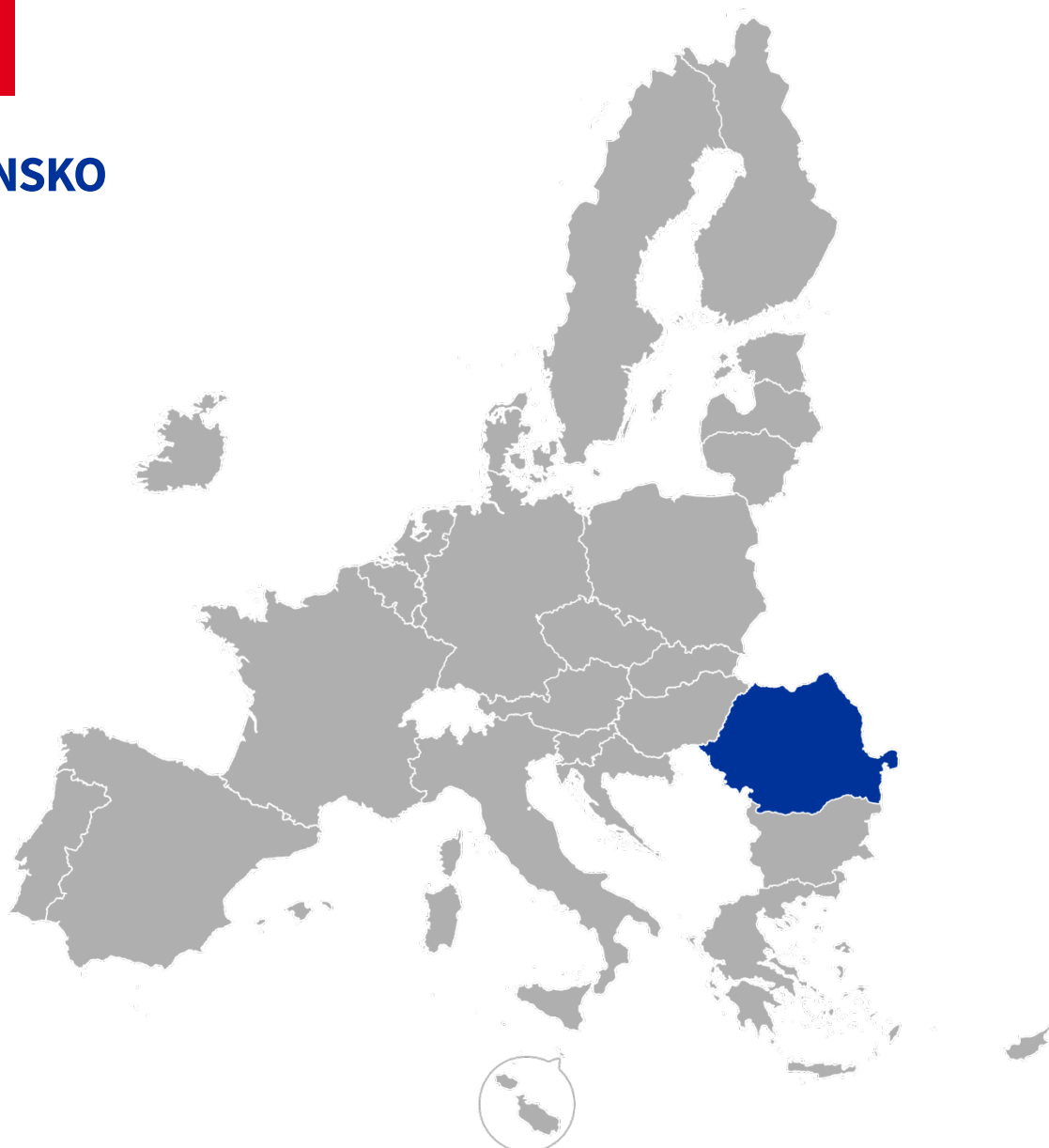
Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



RUMUNSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílají země, jako je Dánsko, Francie, Polsko a Rakousko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vaším cílem je zlepšit život občanů vaší země a zajistit, aby vývoj AI měl co největší pozitivní dopad na společnost.
- Zemědělství je klíčovým odvětvím, v němž může AI výrazně oživit hospodářství vaší země, a proto v zájmu dosažení tohoto cíle podporujete místní rozvoj AI. Upřednostňujete přístup založený na odvětvích, v rámci něhož se na klíčová odvětví, jako je zemědělství, nebudou vztahovat přísnější pravidla, zároveň se však v citlivějších oblastech uplatní přísnější protokoly.
- Požadavky na vývoj „vysoce rizikových“ systémů AI však musí být i nadále přiměřené a je třeba zajistit, aby rámec poskytoval právní jistotu, aniž by vytvářel zbytečnou zátěž nebo náklady pro podniky a spotřebitele. Podporujete proto určité výjimky z požadavků pro vývojáře.
- Zdůrazňujete, že je důležité, aby rámec pro AI podporoval všechny členské státy EU, nejen bohaté země, jako je Německo nebo Francie. Pokud se dnešní jednání zastaví na mrtvém bodě a ukáže se, že kompromisu nelze dosáhnout, jste připraveni souhlasit s přístupem založeným na schopnostech. V tomto případě jste však pro to, aby bylo jasně uvedeno, že zemědělství je odvětvím s nízkým rizikem, aby se v tomto odvětví zabránilo přísným požadavkům na testování.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Zemědělství je sice odvětví, v němž očekáváte největší přínosy AI, ale předpokládáte také významné výhody v oblasti vymáhání práva. Obáváte se však, že pokud by se na donucovací orgány vztahovaly byt minimální požadavky na testování, mohlo by to vést ke zpřístupnění citlivých algoritmů nebo provozních podrobností regulačním orgánům EU, což by mohlo vést k možným únikům informací a ohrozit probíhající vyšetřování.
- Proto se zasazujete o to, aby se na aplikace využívané pro účely národní bezpečnosti nařízení nevztahovalo, a zdůrazňujete, že policie a národní bezpečnostní agentury musí být schopny rychle a pružně zasahovat proti trestné činnosti a terorismu.
- Pokud váš požadavek nebude v dnešní diskusi plně přijat, jste ochotni přistoupit na kompromis, aby nařízení bylo přezkoumáno a případně za dva roky revidováno.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

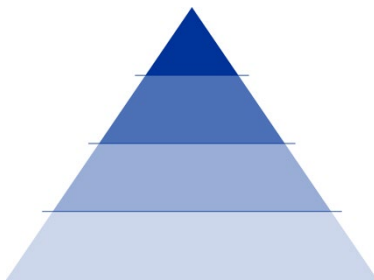
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabíjet bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

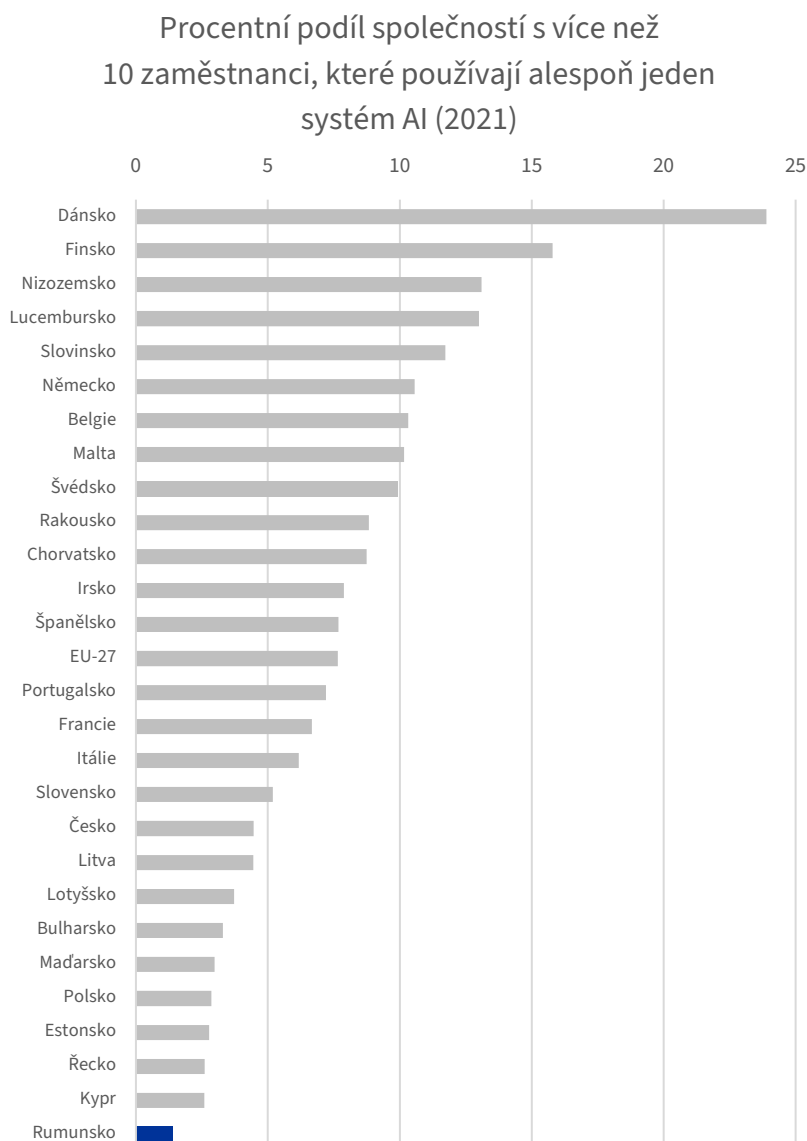
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko	na základě účelu	flexibilní		národní bezpečnost	
Slovensko					
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

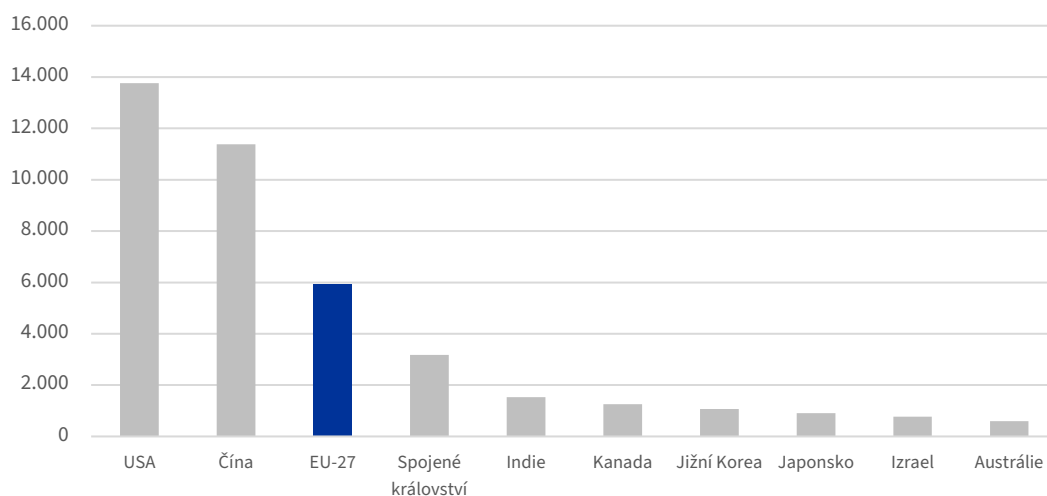
	Evropská unie	Rumunsko
Počet obyvatel	447 695 350	19 054 548
HDP na obyvatele (v EUR)	38 150	17 010
Míra nezaměstnanosti	6,1 %	5,6 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	1,5 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	9 %



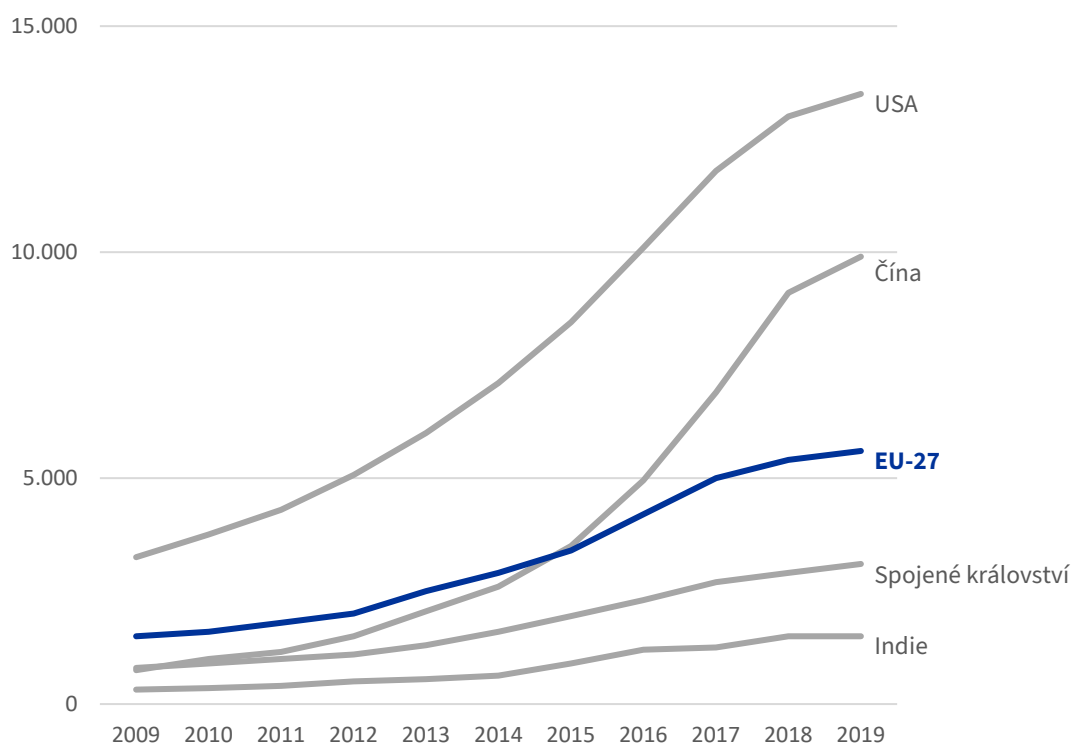
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



SLOVINSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.



Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílají země, jako je Chorvatsko, Itálie a Nizozemsko.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Slovinsko usiluje o komplexní a efektivní digitální transformaci, která bude postavena na celoživotním učení, mezigenerační spolupráci, ale i na zvyšování počtu digitálních odborníků v této zemi. Technologii AI vnímáte jako nástroj, který má být ku prospěchu lidem.
- Podporujete přístup založený na schopnostech, protože lépe odráží skutečná rizika, zejména rizika v oblasti ochrany soukromí. Podle vás je logické klasifikovat všechny technologie AI, které jsou schopny „zpracovávat, analyzovat a vytvářet osobní údaje“ jako „vysoce rizikové“, jak je v přístupu založeném na schopnostech navrženo.
- Během dnešní diskuse chcete upozornit na potenciální rizika, které AI přináší v oblasti lidských práv, spravedlnosti a lidské činnosti. Rizikovost AI se významně zvyšuje, pokud lidé nevědí, jak tato technologie funguje, co dokáže a jaká jsou její omezení. Bez základní gramotnosti v oblasti AI mohou lidé automatizovaným systémům slepě důvěřovat a nemusí být schopni rozeznat předpojatost nebo manipulace. V konečném důsledku pak ani nejsou schopni vymáhat odpovědnost vývojářů nebo institucí. Nedostatečné znalosti v oblasti AI také zvyšují zranitelnost vůči neúmyslně nesprávným informacím a prohlubují sociální nerovnosti. Také mohou být příčinou iracionálního strachu z AI, nebo mohou naopak vést k jejímu nekritickému využívání. Aby tedy AI mohla být používána bezpečně a bez obav, je nezbytná rozsáhlá osvěta.
- Pokud bude v aktu o umělé inteligenci význam osvěty veřejnosti dostatečně zdůrazněn, jste otevřeni tomu, aby byla při vývoji AI umožněna větší flexibilita.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Podle vašeho názoru by akt EU o umělé inteligenci měl výrazně urychlit rozvoj AI ve všech členských státech EU. Zavedené společnosti působící v oblasti AI sídlí v několika zemích s velkým kapitálem. Oproti tomu jiné členské státy jsou stále ve fázi, kdy vytvářejí příznivé podmínky pro startupy. Podstatou startupů je vytvořit v krátkém čase a s omezeným přístupem ke kapitálu něco převratného. To je při velké administrativní zátěži téměř nemožné.
- Pokud bude na základě článku 1 zaveden rámec pro předběžné testování (což podporujete), navrhuje, aby se požadavek na testování nevztahoval na startupy s méně než 20 zaměstnanci a ročním obratem do 4 milionů eur. Podle vás by se měl u startupů uplatňovat spíše flexibilní přístup. Považujete to za vyvážený kompromis a chcete o něm vést diskusi.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

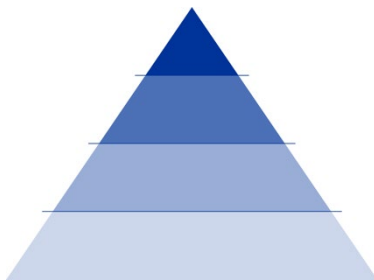
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabíjet bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko	na základě schopností	předběžná opatrnost		startupy	
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

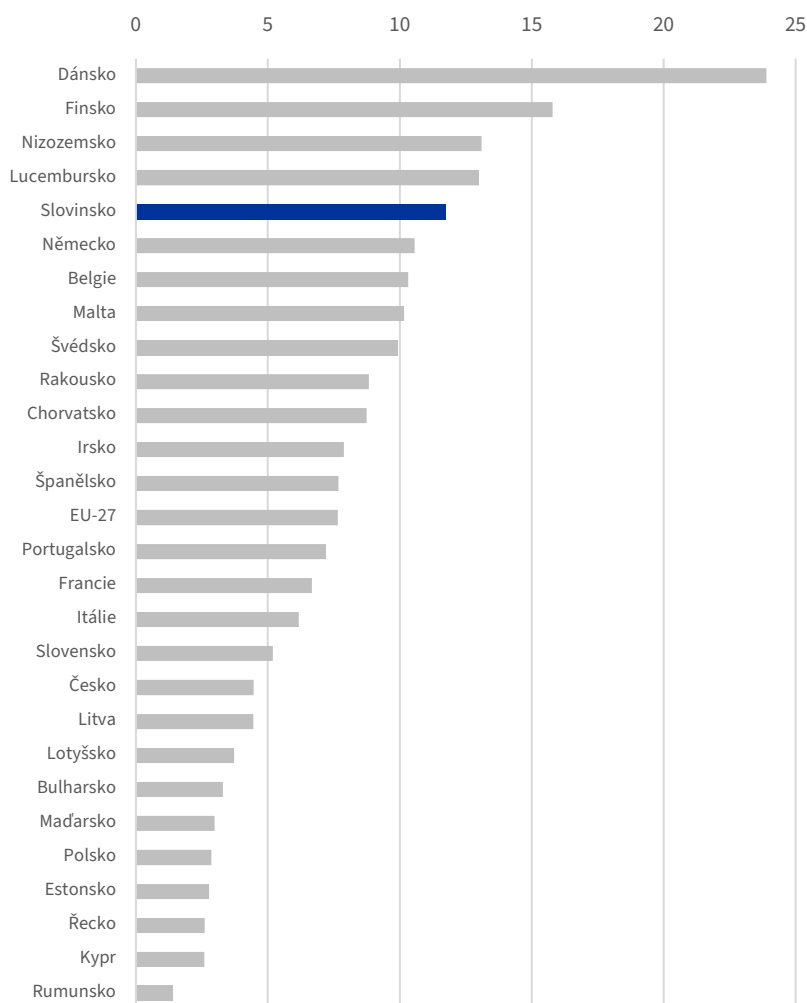
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Slovinsko
Počet obyvatel	447 695 350	2 116 972
HDP na obyvatele (v EUR)	38 150	30 160
Míra nezaměstnanosti	6,1 %	3,7 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	11,4 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	18,9 %

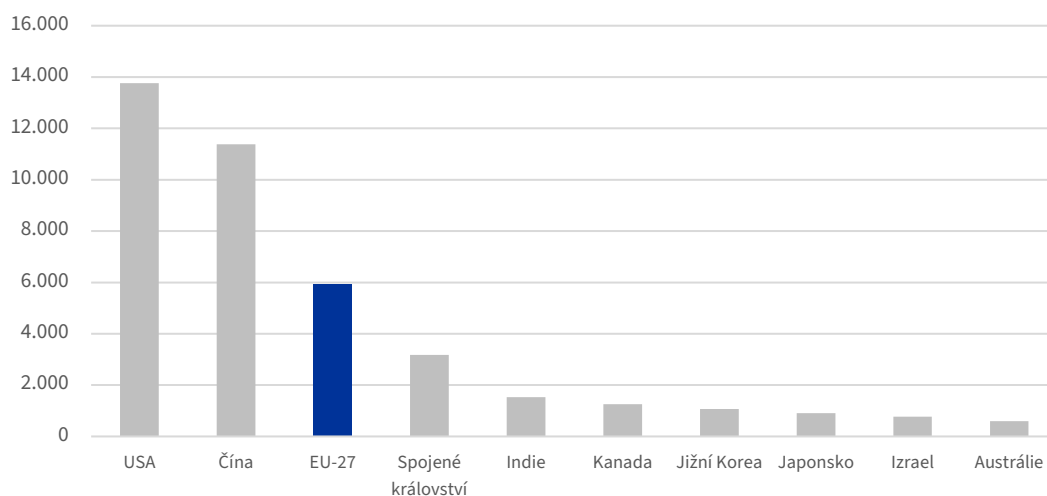
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



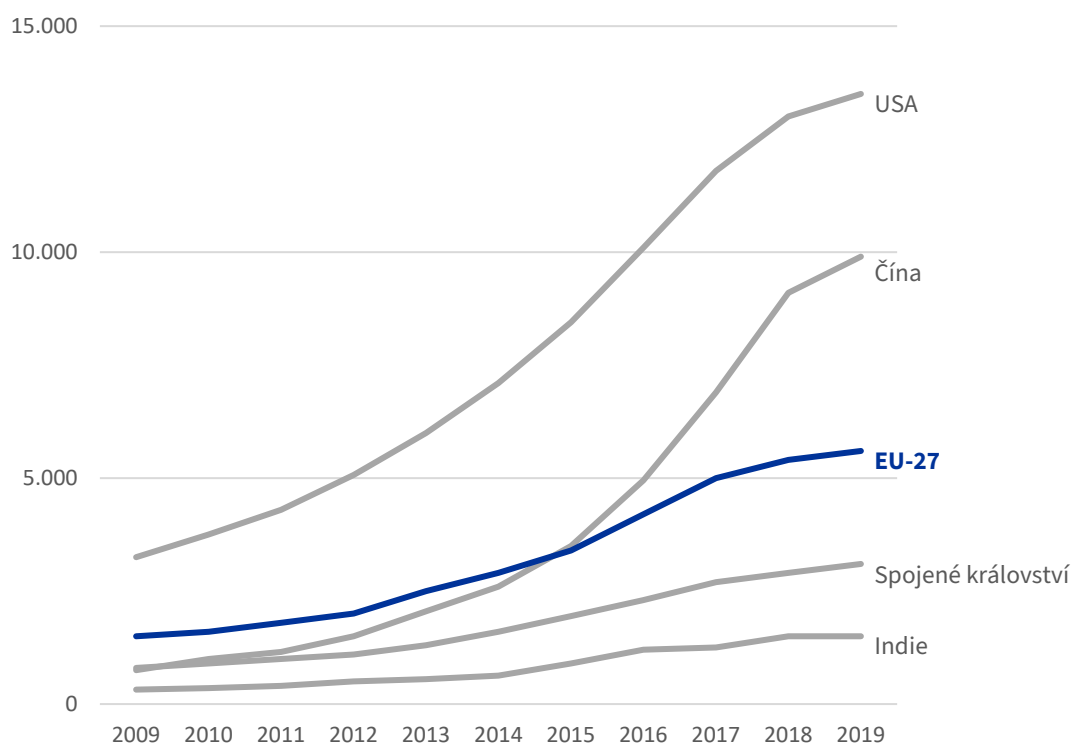
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



SLOVENSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **přečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílají země, jako je Bulharsko, Lucembursko a Malta.** Jednejte se se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost **do tabulky na stranách 8 a 9.** Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty.** Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás.** Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Slovensko chce najít způsob, jak nejlépe využít akt EU o umělé inteligenci. Dosud nebyly na vnitrostátní výzkum AI vyčleněny žádné veřejné finanční prostředky. Vaše vláda se proto spoléhá na soukromé investice nebo podporu ze strany jiných členských států EU.
- Podporujete přístup založený na účelu, protože se jedná o způsob, jakým EU obvykle reguluje technologie a produkty, např. zdravotnické prostředky nebo chemické výrobky. Jste však toho názoru, že současný seznam vysoce rizikových odvětví by měl být omezen na kritickou infrastrukturu a oblast zdraví, neboť se jedná o odvětví, která mají přímý dopad na záležitosti týkající se života a smrti. Další odvětví, přestože jsou důležitá, mohou být přehodnocena později. Po konzultaci s občanskou společností a průmyslem jste ochotni tento seznam za tři roky přezkoumat.
- Pokud jde o požadavky na testování, obáváte se, že příliš striktní pravidla by mohla vyvolat odliv mozků a způsobit, že evropské talenty v oblasti AI se přesunou do flexibilnějšího prostředí, jako jsou USA. Pro Slovensko jako pro jeden z menších členských států EU, které mají k dispozici jen málo zdrojů na investice do AI, může být přínosný flexibilní přístup.
- Pokud dnešní jednání skončí na mrtvém bodě a bude jasné, že kompromisu nelze dosáhnout, jste ochotni přijmout lehce přísnější požadavky na testování. V rámci vyváženého kompromisu by však měla být **méně bohatým členským státům EU poskytnuta finanční podpora**, která by pomohla pokrýt zvýšené náklady na rozvoj.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- EU dosud ve vývoji AI celosvětově zaostává. Výjimka pro startupy by pomohla vytvořit dynamický ekosystém. Startupům nemají právní, finanční a správní kapacity, aby mohly dodržovat stejná pravidla jako velké technologické společnosti. Přísné předpisy by mohly zbrzdit inovace hned v jejich počáteční fázi.
- Případně, pokud se vám nepodaří přesvědčit dostatečný počet členských států, aby přijaly flexibilní požadavky na testování uvedené v článku 1, chcete zajistit, aby startupům byly uděleny určité výjimky z těchto požadavků. Nebudou tak hned na začátku čelit zátěži spojené s úplným dodržováním předpisů (pokud to vyplyne z článku 1), ale budou spíše schopny fungovat v experimentálnějších prostředí.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

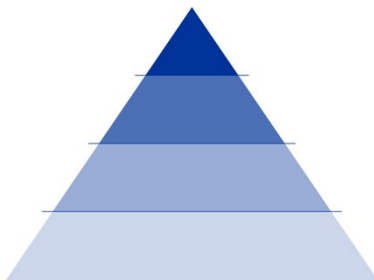
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabíjet bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko	na základě účelu	flexibilní	finanční podpora	startupy	
Slovinsko					
Španělsko					
Švédsko					
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

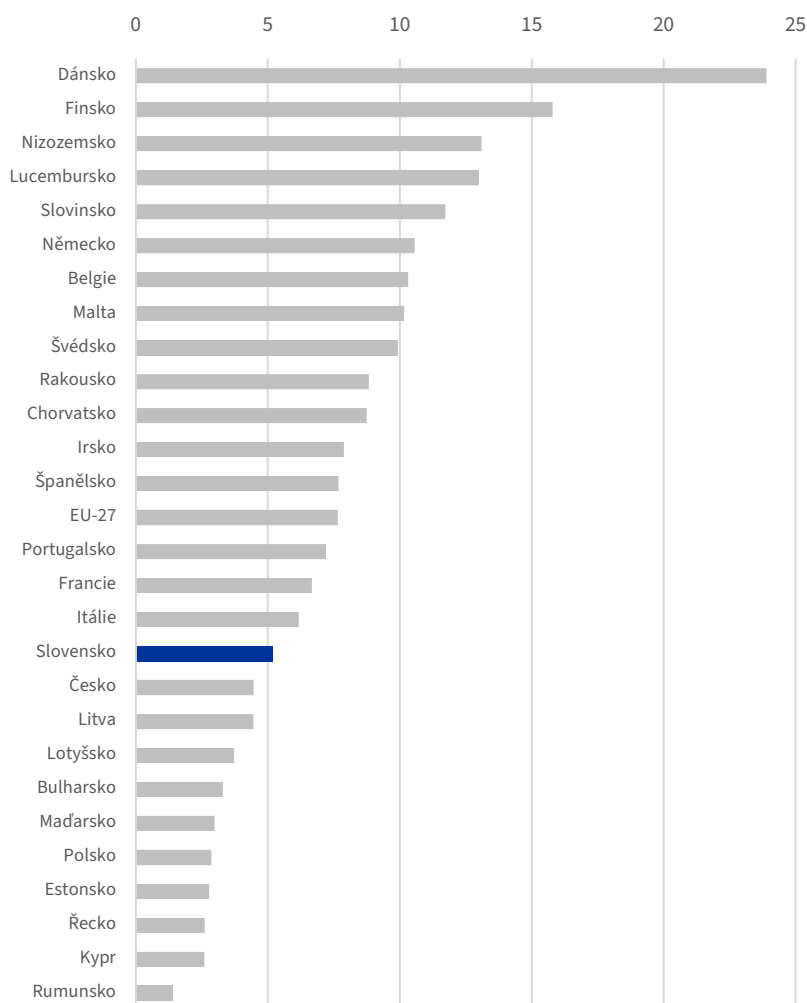
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

	Evropská unie	Slovensko
Počet obyvatel	447 695 350	5 428 792
HDP na obyvatele (v EUR)	38 150	22 690
Míra nezaměstnanosti	6,1 %	5,8 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	7 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	21,7 %

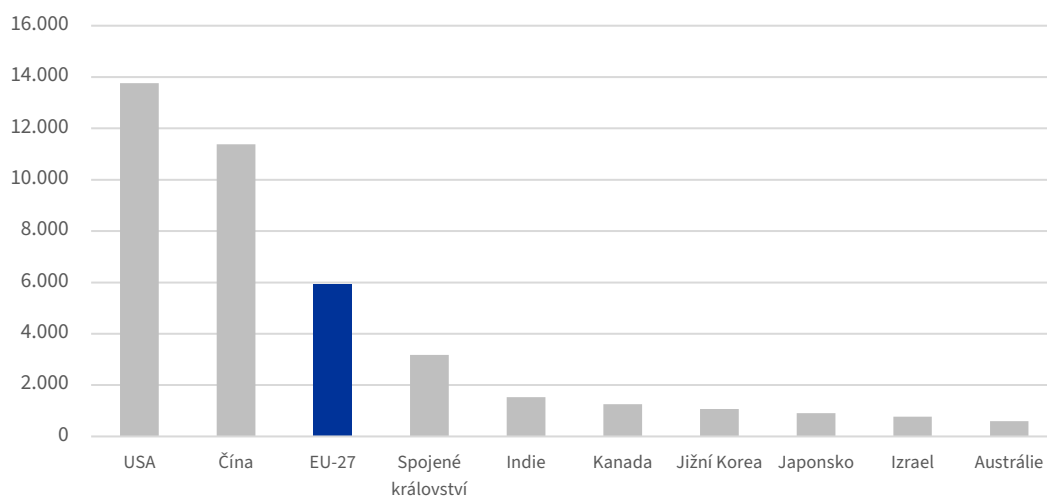
Procentní podíl společností s více než
10 zaměstnanci, které používají alespoň jeden
systém AI (2021)



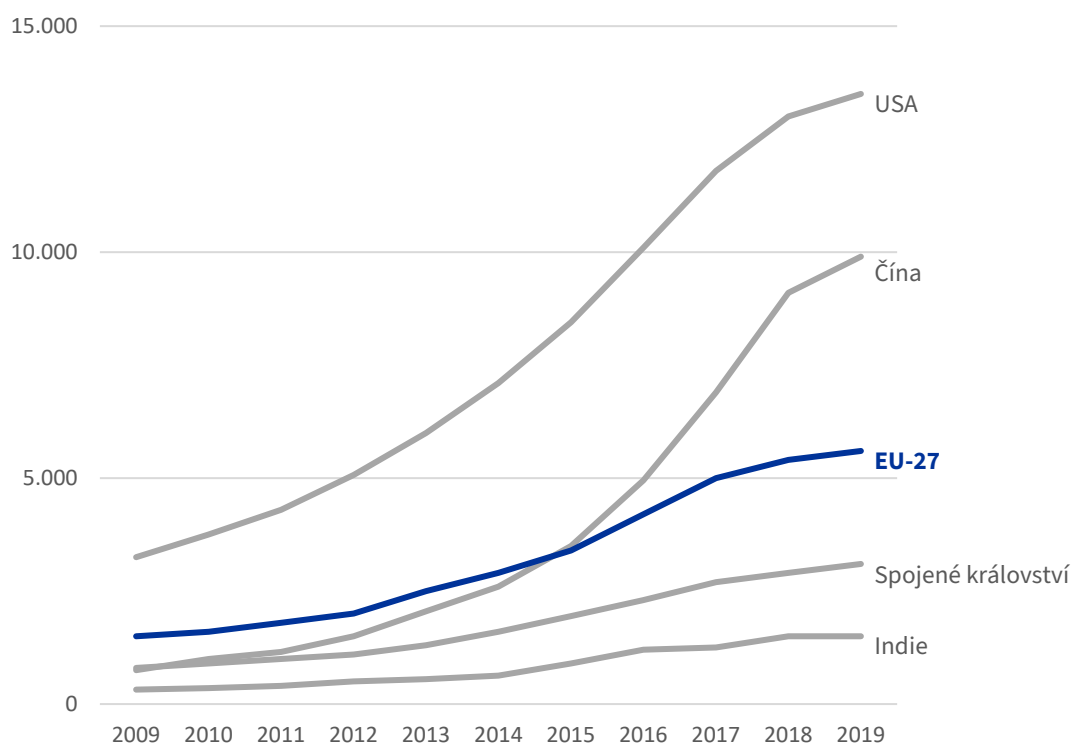
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

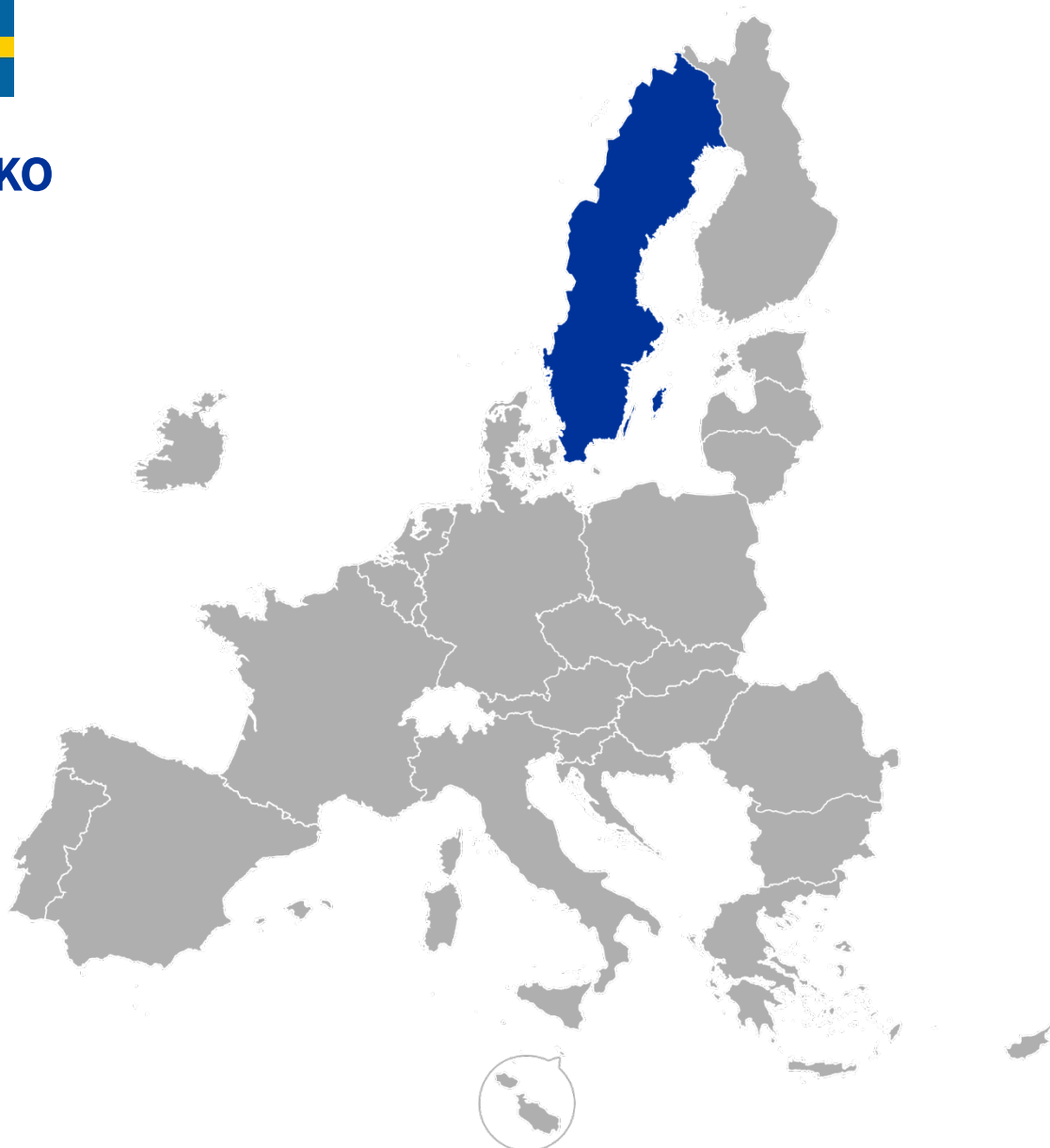
Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.
© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print	ISBN 978-92-850-0496-5	doi: 10.2860/6030812	QC-01-25-006-CS-C
PDF	ISBN 978-92-850-0495-8	doi: 10.2860/8915465	QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



ŠVÉDSKO



NA DEN MINISTREM: JEDNÁNÍ NA ÚROVNI EU-27

■ Zavedení rámce pro AI

Popisy rolí | **Verze pro pokročilé** | CS



Rada
Evropské unie

Návrh

NAŘÍZENÍ EVROPSKÉHO PARLAMENTU A RADY

o právním rámci pro využívání umělé inteligence

KONTEXT NÁVRHU

Evropská unie se nachází v kritickém bodu utváření své digitální budoucnosti. V rámci Strategie kybernetické bezpečnosti pro digitální dekádu si EU klade za cíl stát se lídrem v oblasti umělé inteligence a zároveň zajišťovat základní práva, ochranu spotřebitele, konkurenceschopnost trhu, inovace a vedoucí postavení v oblasti technologií. Jádrem diskusí je legislativní návrh Evropské komise, jehož cílem je regulovat vývoj umělé inteligence v EU. Proběhlo již několik kol jednání. Je ale třeba ještě vyřešit následující otázky.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

- A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Při definování „vysoce rizikových“ systémů AI přijme EU **klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.**
- B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat **přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování** (viz příloha II).

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI s **výjimkou těch, které jsou vyvíjeny pro vojenské nebo obranné účely.**

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



ŠVÉDSKO

Pořad jednání

Dnes je na pořadu jednání Rady návrh Komise týkající se právního rámce pro využívání umělé inteligence. O návrhu budete jednat se zástupci ostatních členských států tak, abyste dosáhli kompromisu.

Váš hlavní cíl

Společný postoj Rady zohledňující vaše zájmy. Můžete rovněž podpořit kompromisní návrh, který vaše zájmy zohledňuje částečně.

Váš úkol

1. Pečlivě si **prečtěte informace týkající se** zájmů a postojů **vaší skupiny**.
2. **Některé z vašich argumentů s vámi sdílají země, jako je Finsko, Irsko a Portugalsko.** Jednejte se zástupci těchto zemí a zjistěte, zda s nimi můžete spolupracovat na prosazování případných společných zájmů.
3. Připravte si **úvodní prohlášení** pro první kolo oficiálního jednání (prohlášení by mělo trvat max. 1 minutu). Uvedte svůj postoj k článkům 1 a 2. Může vám pomoci text uvedený v šedém rámečku vpravo.
4. Úvodní prohlášení přednesete poté, co vám předsednictví (jehož zástupce zasedání řídí) udělí slovo. Pokud souhlasíte s postoji jiných zemí, které svůj příspěvek již přednesly, uveďte to.
5. Vyslechněte si pozorně postoje ostatních zemí. **Poznamenejte si je** pro přehlednost do **tabulky na stranách 8 a 9**. Vyjednávání z velké části spočívá v naslouchání ostatním účastníkům a v hledání řešení, která budou přijatelná pro všechny.
6. Během oficiálního jednání **předložte své postoje a argumenty**. Pokud se vaše argumenty shodují s argumenty jiných zemí, může být výhodné je přednést společně, aby jim byla přikládána větší váha.
7. Ovšem **to, jak budete nakonec hlasovat, je samozřejmě pouze na vás**. Pamatujte na to, že při závěrečném hlasování o kompromisním návrhu se musíte v podstatě rozhodnout, zda upřednostňujete kompromis, se kterým možná nebudete stoprocentně spokojeni, nebo zda dáváte přednost tomu, aby jednání skončilo zcela bez výsledku.

*Paní předsedkyně / pane předsedo, paní komisařko /
pane komisaři, kolegyně a kolegové,*

*nejprve bych chtěl/a poděkovat Komisi za
vypracování návrhu a předsednictví za jeho zařazení
na pořad dnešního jednání.*

Před dnešním zasedáním jsme vedli plodné diskuse s

*Domníváme se, že je nezbytné přijmout klasifikaci rizik
založenou na schopnostech/účelu, protože*

*Pokud jde o článek 1B, domníváme se, že během
vývoje „vysoce rizikových“ systémů AI musí vývojáři
uplatňovat flexibilní přístup / přístup předběžné
opatrnosti, protože*

*Pokud jde o článek 2, některé výjimky budou/nebudou
platit, protože jsme pro*

*Věříme, že se nám dnes podaří nalézt přijatelný
kompromis.*

Děkuji vám za pozornost.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 1

Klasifikace systémů umělé inteligence a řízení rizik

A) EU přijímá společnou normu pro řízení rizik při vývoji, uvádění na trh, dovozu a používání systémů umělé inteligence (AI). Tyto systémy jsou klasifikovány jako systémy obnášející nulové, nízké, vysoké nebo nepřijatelné riziko (viz příloha I). Pro účely definice „vysoce rizikových“ systémů AI přijme EU



...klasifikaci rizik založenou na schopnostech: riziko se stanoví se zřetelem k technickým schopnostem systému.



...klasifikaci rizik založenou na účelu: riziko se stanoví na základě zamýšleného použití AI v citlivých odvětvích.

Na základě této klasifikace rizik musí „vysoce rizikové“ systémy AI před zprovozněním projít přísným testováním.

B) Během vývoje „vysoce rizikových“ systémů AI musí vývojáři uplatňovat



...flexibilní přístup: před zprovozněním modelu AI provést vlastní hodnocení.



...přístup předběžné opatrnosti: dodržovat přísné standardizované požadavky na testování (viz příloha II).

Vaše argumenty

- Vláda vaší země podporuje progresivní, ale praktickou regulaci AI. V návrhu prosazujete jen omezené výjimky, aby nebyli odrazeni investoři a výzkumní pracovníci v oblasti AI.
- Samotné Švédsko nemůže v oblasti AI hrát vedoucí úlohu, avšak jednotná EU to dokázat může – tím, že bude podporovat inovace, stanovovat normy a posilovat svůj globální vliv. K dosažení tohoto cíle je zapotřebí více motivovat vývojáře AI a omezit regulační překážky ve fázích vývoje, aby se podpořila tvorba pracovních míst a hospodářský růst.
- Klíčový je přístup založený na účelu: AI používaná v lékařské diagnostice by měla podléhat přísnějším předpisům než umělá inteligence určená pro řízení podniků. Nechcete zlehčovat obavy týkající se například předpojatosti AI, ale jde vám o to, aby byla pozornost přednostně věnována rizikům, která jsou skutečně bezprostřední – těm, která přímo ovlivňují lidské životy. Stanovení rizika s ohledem na odvětví navíc znamená menší počet potenciálních konfliktů se stávajícími odvětvovými požadavky obsaženými v jiných předpisech, jako jsou například rámce upravující zdravotní péči a finance. Administrativní složitost je tak udržována na zvládnutelné úrovni.
- Pokud jde o požadavky na vývojáře, pevně věříte, že evropské technologické společnosti mohou vyvinout vysoce kvalitní AI bez přísné regulace. Pokud se vývojářům poskytne flexibilita při testování a zdokonalování jejich modelů, podpoří se tím inovace a excelence a zároveň zajistí odpovědný vývoj AI.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Článek 2

Oblast působnosti

Ustanovení uvedené v článku 1 se vztahuje na všechny systémy AI



bez výjimky



bez výjimky, není-li systém vyvinut jako systém s otevřeným zdrojovým kódem



s výjimkou případů, kdy jsou vyvinuty malými podniky a startupy (a)



s výjimkou případů, kdy jsou vyvinuty pro vojenské nebo obranné účely (a)



s výjimkou případů, kdy jsou vyvinuty pro účely národní bezpečnosti, řízení migrace nebo ochrany hranic

Vaše argumenty

- Má pro vás zásadní význam, aby do oblasti působnosti nařízení o AI nespadaly záležitosti týkající se národní bezpečnosti, armády nebo obrany. Členské státy jsou odhodlány zachovat si kontrolu nad rozhodnutími v těchto oblastech, které stejně za normálních okolností nespadají do pravomoci EU. To se odráží i ve skutečnosti, že Unie nemá jednotnou armádu EU.
- Dále argumentujete tím, že systémy AI používané pro shromažďování zpravodajských informací, kybernetickou obranu nebo plány vojenských operací nemohou podléhat stejným požadavkům na zveřejnění nebo regulačním požadavkům jako civilní systémy, protože by tak mohlo být ohroženo utajení operací, zájmy národní bezpečnosti a strategická výhoda, kterou mají tyto systémy poskytovat.
- Článek 2 je pro vás důležitější než článek 1. Pokud se vám nepodaří přesvědčit ostatní o svých zájmech týkajících se článku 1, zaměřte se na to, abyste prosadili vaše zásadní požadavky vztahující se k článku 2.
- Vaše poznámky:

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

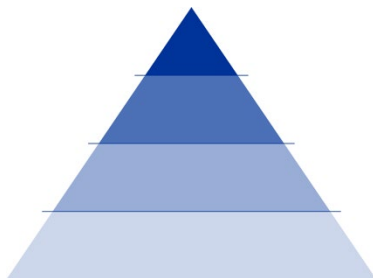
Příloha I: Klasifikace rizik

Nepřijatelné riziko

Vysoké riziko

Nízké riziko

Nulové riziko



Úplný zákaz

Přísné požadavky (viz příloha II)

Požadavky na transparentnost

Žádná regulace

Nepřijatelné riziko

- **AI pro hodnocení sociálního kreditu** – seřazení občanů do kategorií podle jejich chování
- **Hromadné sledování pomocí AI** – biometrické sledování bez souhlasu
- **Rozpoznávání emocí na pracovištích / ve školách** s cílem ovlivnit dosahované výsledky
- **Autonomní smrtící zbraně** – schopné zabít bez lidské intervence
- **Podprahová manipulace pomocí AI** – netransparentní ovlivňování lidí při rozhodování

Vysoké riziko



Přístup založený na schopnostech

Všechny systémy GPAI a AI, které jsou schopny:

- přijímat zásadní rozhodnutí o právech, bezpečnosti nebo živobytí lidí – např. umělá inteligence pomáhající v chirurgii, automatizovaný nábor zaměstnanců
- samostatně ovlivňovat nebo manipulovat lidské chování v citlivých oblastech – např. mikrocílení politické reklamy
- zpracovávat, analyzovat nebo vytvářet osobní/biometrické údaje ve velkém měřítku
- provádět složité prognózy nebo posouzení rizik se společenskými důsledky – např. klimatické modely
- pracovat v prostředí, kde by chyby nebo zkreslení mohly způsobit újmu, např. samořídící automobily

Přístup založený na účelu

Všechny systémy AI určené k použití:

- v kritické infrastruktuře
- ve zdravotnictví
- na pracovišti (HR)
- v souvislosti s dávkami sociálního zabezpečení
- v rámci justičních orgánů

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Nízké riziko

- **Systémy AI neklasifikované jako „vysoce rizikové“**
- **Chatboty a virtuální asistenti** (pokud uživatelé vědí, že hovoří s AI)
- **Doporučovací algoritmy** – návrhy s využitím AI na sociálních sítích
- **AI pro automatizaci podnikání** – AI, která automatizuje základní administrativní úkoly

Nulové riziko

- Generátory/editory využívající AI ke zpracování či tvorbě obrazových, hudebních či jiných uměleckých materiálů
- Umělá inteligence ve vědeckém výzkumu
- Umělá inteligence v hračkách a zábavních robotech – hračky vybavené umělou inteligencí bez rizik v reálném světě



Příloha II: Požadavky na „vysoce rizikové“ systémy AI

Důraz na právní a etické normy:

- Posouzení rizik před uvedením systému do provozu, identifikace předvídatelných rizik a možného zneužití.
- Zajištění lidského dohledu, aby se zabránilo plné automatizaci při rozhodování.
- Správa dat k zajištění objektivních trénovacích dat.
- Povinné externí audity, včetně testování na pískovišti s lidským dohledem.
- Zpřístupnění výsledků testů regulačním orgánům EU.
- Monitorování výsledků po uvedení na trh.

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.



Definice

- **Algoritmy** – pokyny, kterými se počítač řídí krok za krokem za účelem vyřešení problému nebo přijetí rozhodnutí
- **Systémy umělé inteligence** (systémy AI) – počítačové systémy, které využívají umělou inteligenci k plnění úkolů, jež obvykle vyžadují lidskou inteligenci, např. učení, uvažování, řešení problémů, rozhodování, vnímání a porozumění jazyku
- **Odliv mozků** – když kvalifikovaní pracovníci opustí svou zemi za prací v zahraničí
- **Migrace podniků** – když společnost přesune svou činnost do jiné země, často kvůli lepším podmínkám nebo nižším nákladům
- **AI uplatňující hluboké učení (deep learning)** – systém, který se podobně jako lidský mozek učí z velkého množství dat za využití neuronové sítě. Rozpoznává vzorce a postupem času se zlepšuje, aniž by bylo nutné explicitní programování
- **Umělá inteligence pro všeobecné účely (GPAI)** – umělá inteligence, která může plnit více než jeden úkol, např. ChatGPT, Google Gemini
- **Otevřený zdroj** – software, k jehož kódu má kdokoli volný přístup, může si jej zobrazit, používat a pozměňovat
- **AI založená na pravidlech** – systém AI, který přijímá rozhodnutí na základě předem stanovených pravidel a logiky. Například „Jestliže nastane X, pak proved' úkon Y.“
- **Testování na pískovišti** – testování v izolovaném prostředí, např. umělá inteligence pro autonomní řízení je před jízdou na skutečných silnicích testována v simulaci města

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Rakousko					
Belgie					
Bulharsko					
Chorvatsko					
Kypr					
Česko					
Dánsko					
Estonsko					
Finsko					
Francie					
Německo					
Řecko					
Maďarsko					
Irsko					

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

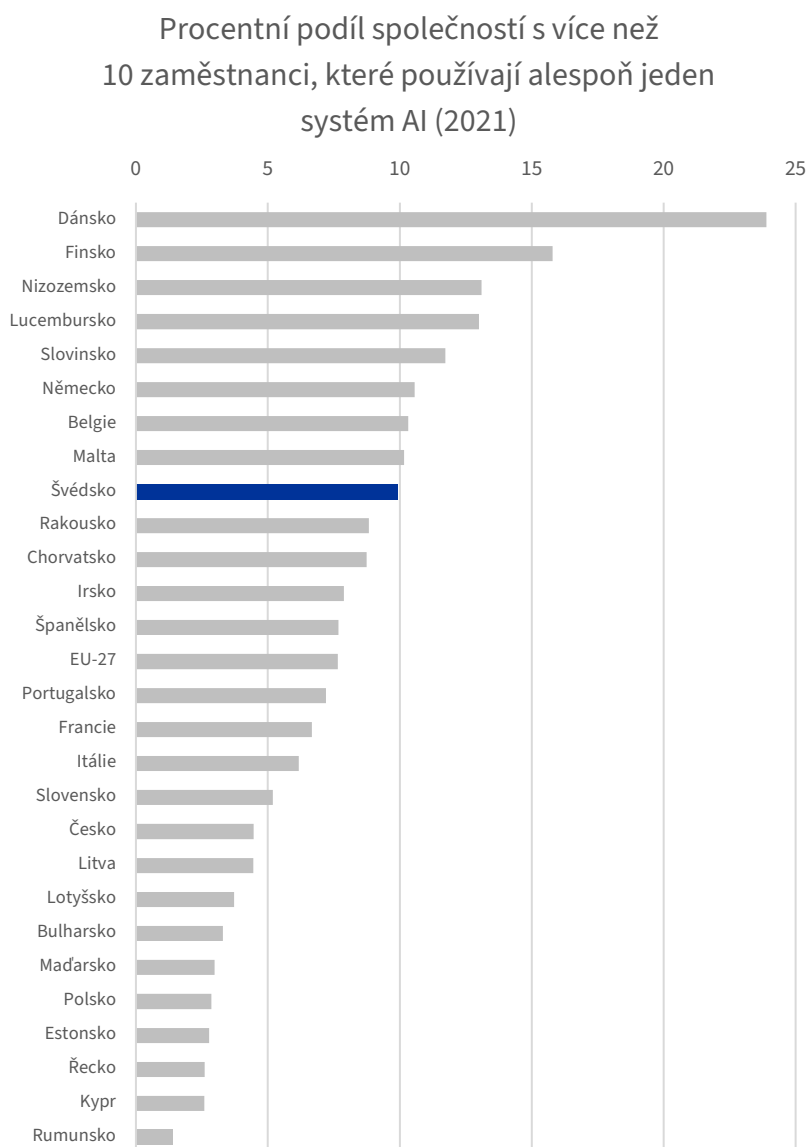
	Článek 1A: klasifikace	Článek 1B: přístup	Další návrhy/poznámky	Článek 2: oblast působnosti	Další návrhy/poznámky
Itálie					
Lotyšsko					
Litva					
Lucembursko					
Malta					
Nizozemsko					
Polsko					
Portugalsko					
Rumunsko					
Slovensko					
Slovinsko					
Španělsko					
Švédsko	na základě účelu	flexibilní		armáda/obrana + národní bezpečnost	
Komise	na základě účelu	předběžná opatrnost		armáda/obrana	

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení. Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Souvislosti

Pozor: **Jednání probíhají v roce 2023!**

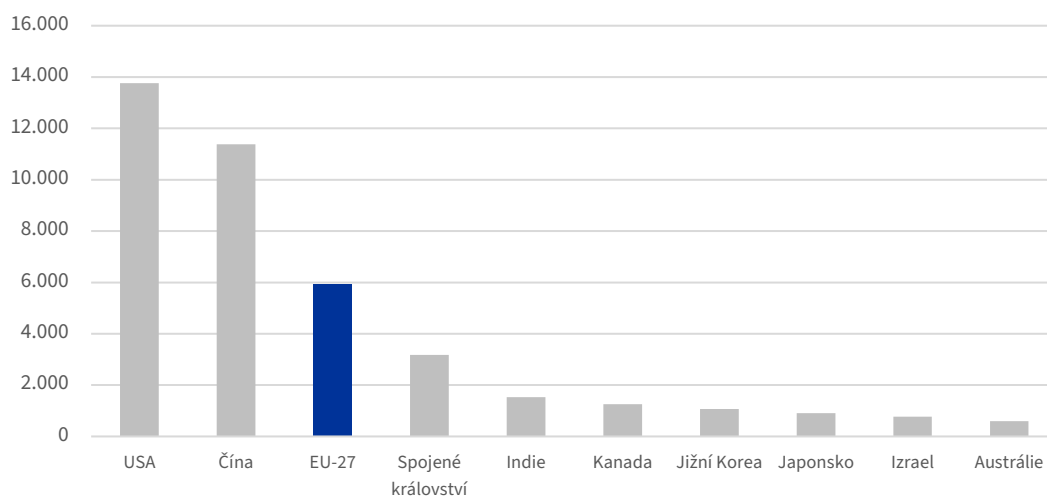
	Evropská unie	Švédsko
Počet obyvatel	447 695 350	10 521 556
HDP na obyvatele (v EUR)	38 150	51 060
Míra nezaměstnanosti	6,1 %	7,7 %
Procentní podíl společností využívajících technologii umělé inteligence	8 %	10,4 %
Podíl obyvatelstva s pokročilými digitálními dovednostmi	27,3 %	36,5 %



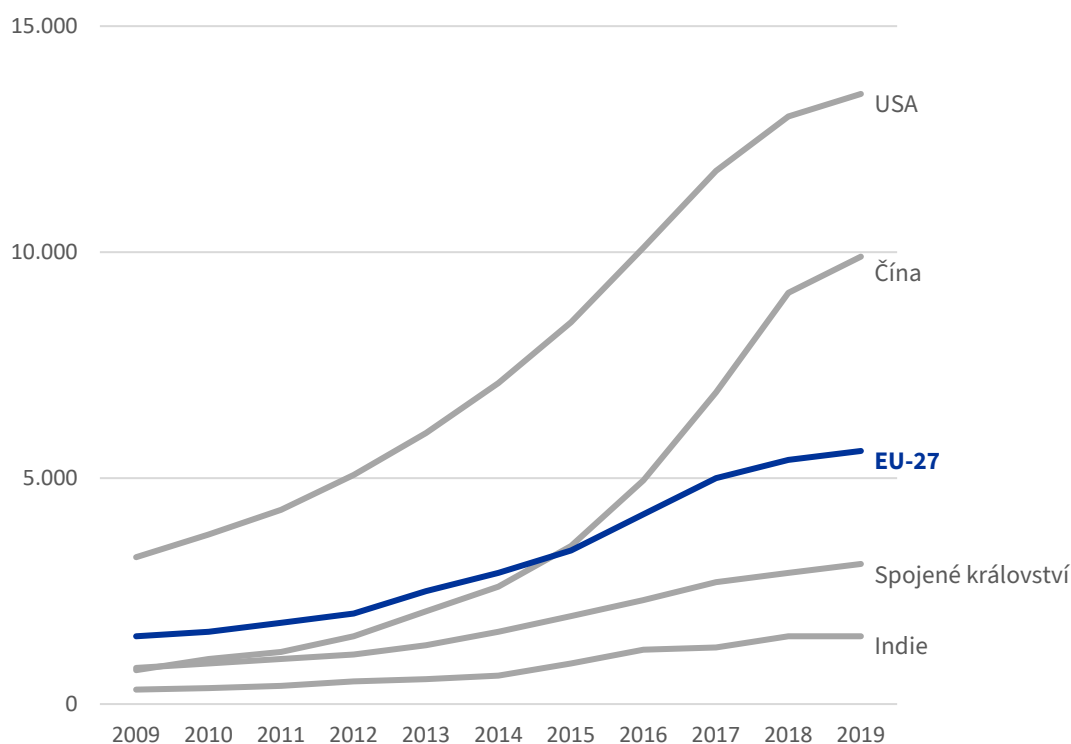
Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.

Společnosti zabývající se AI podle zemí (2020)



Nárůst počtu společností zabývajících se AI



Zdroje údajů: Eurostat, Společné výzkumné středisko

Tuto publikaci vypracoval generální sekretariát Rady, a to pouze pro informační účely. Nezakládá odpovědnost orgánů EU ani členských států. Další informace o Evropské radě a o Radě EU naleznete na internetových stránkách <https://www.consilium.europa.eu/cs/>.

© Evropská unie, 2025 | Reprodukce povolena s uvedením zdroje.

Print ISBN 978-92-850-0496-5 doi: 10.2860/6030812 QC-01-25-006-CS-C
PDF ISBN 978-92-850-0495-8 doi: 10.2860/8915465 QC-01-25-006-CS-N

Poznámka: Tyto materiály byly z didaktického hlediska vypracovány tak, aby byl zajištěn hladký průběh hry i její přínos pro účastníky z hlediska učení.

Jejich obsah nemusí nezbytně odpovídat realitě nebo reálným politickým postojům.