

Jak umělá inteligence mění audit

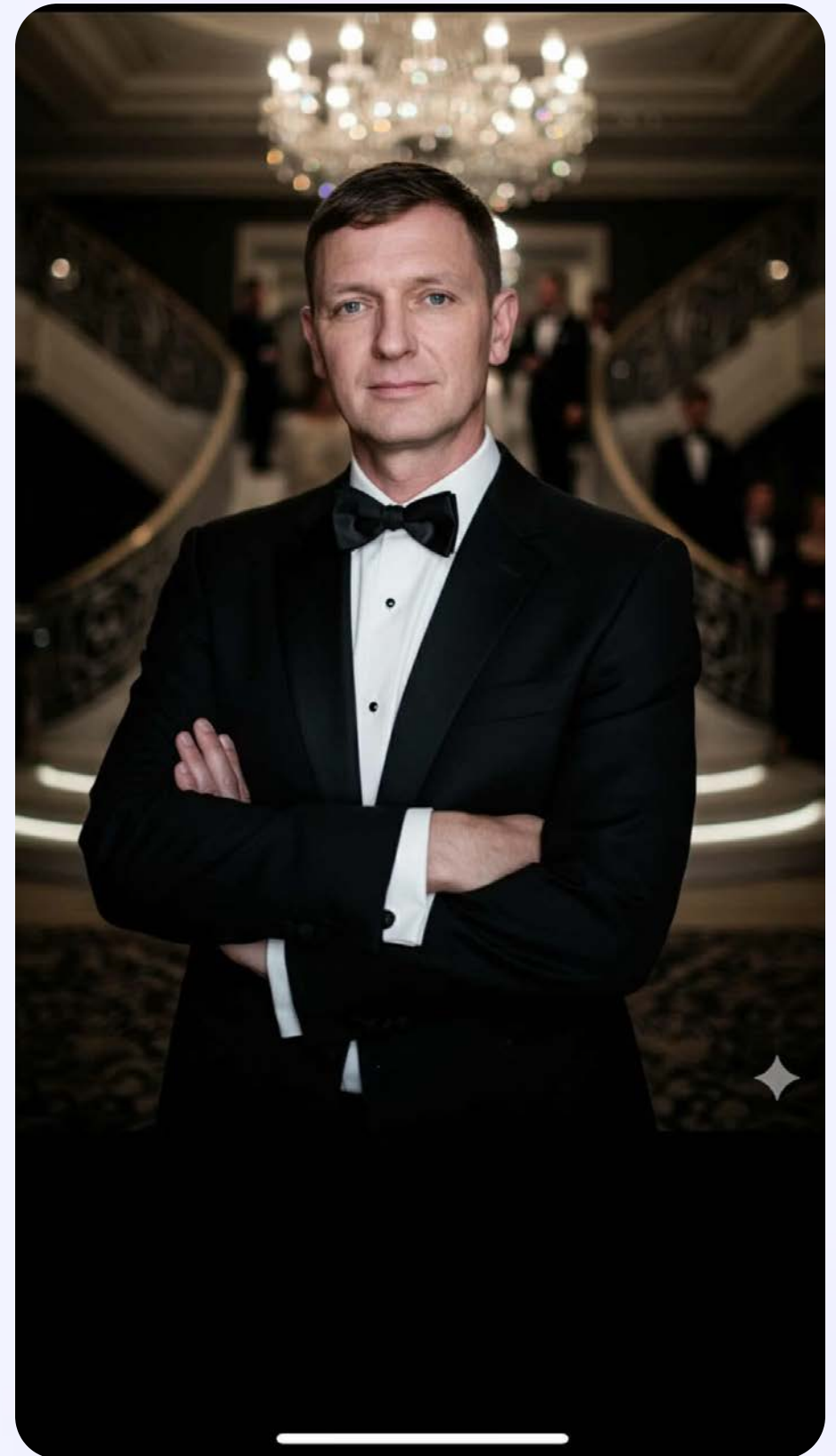
Transformace interního a externího auditu v éře AI probíhá rychleji než čekáme

Jiří Diepolt – listopad 2025



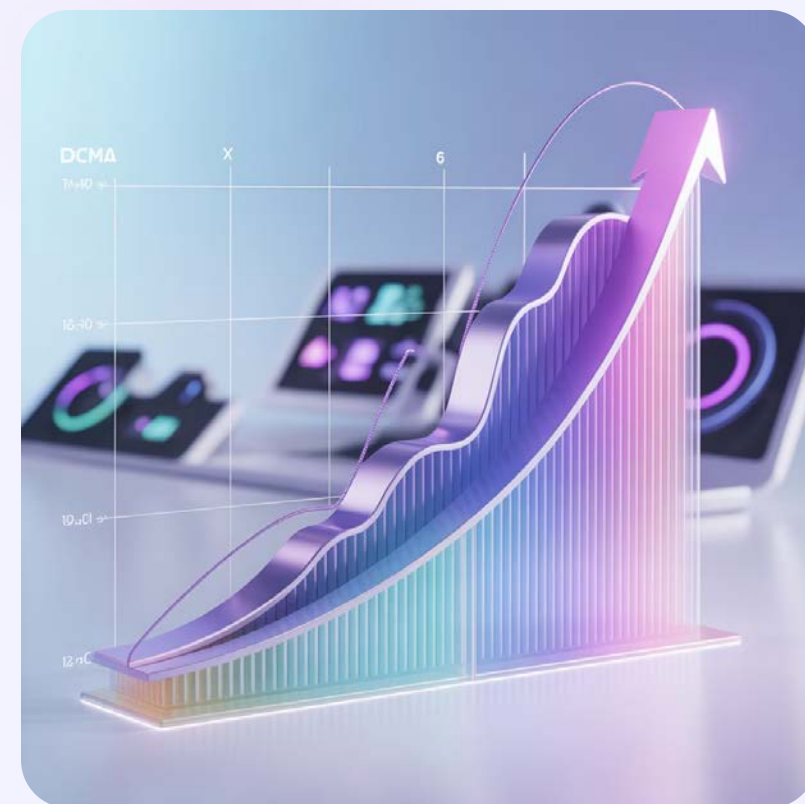
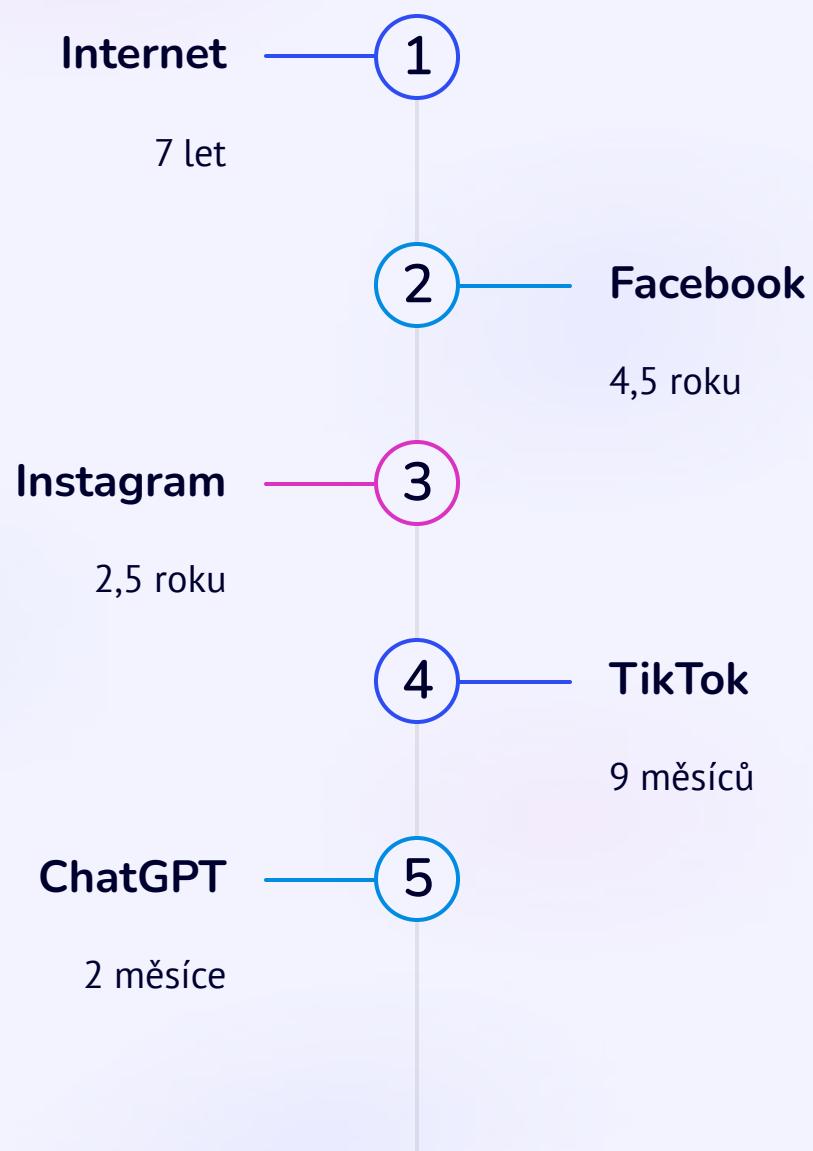
Agenda

- **Kontext a vývoj AI**
Exponenciální růst adopce umělé inteligence a její dopad na auditní profesi
- **Praktické využití v auditu**
Konkrétní aplikace AI v jednotlivých fázích auditního procesu
- **Rizika a governance**
Identifikace rizik, reálné incidenty a principy řízení AI
- **Bezpečné nasazení**
Sedm klíčových opatření pro minimalizaci rizik při využití AI
- **Klíčové otázky**
Praktické otázky pro auditory partnery a doporučené informační zdroje
- **Shrnutí**



Exponenciální růst adopce AI

Čas potřebný k dosažení **100 milionů** uživatelů:



- ❏ **Závěr:** AI není budoucnost – je přítomnost. Rychlost adopce AI technologií překonává vše, co jsme v historii technologií (zatím) zažili.

Výsledky průzkumu řízení rizik AI v ČR

87%

Organizací již AI využívá

Většina firem v ČR už nějakou formu AI nasadila

21%

AI rizika začleněna v systému řízení rizik

Pouze 21% má rizika AI začleněna do systému řízení rizik

41%

Žádný formální rámec

41% organizací nepoužívá žádný framework pro řízení rizik AI

34%

Ad hoc přístup

34% organizací řídí rizika AI ad hoc, bez systematického přístupu

Podle průzkumu České asociace umělé inteligence (ČAAI) z roku 2025 jsou klíčové výzvy v řízení rizik AI v českých organizacích:

- **Rychlý vývoj AI technologií** vnímají jako největší bariéru (52%)
- **Obtížnou měřitelnost rizik** souvisejících s AI uvedlo 34% dotázaných
- **Konflikt mezi potřebou inovovat a udržet kontrolu** řeší 33% organizací



Evoluce využití AI v auditu

Fáze transformace auditní profese z pohledu AI



Pět hlavních bariér zavádění AI



Maturity Model AI

Kde se nachází vaše organizace na stupnici zralosti využívání umělé inteligence?



Úroveň 1: Ad-hoc

Nekoordinované experimentování s veřejnými AI nástroji, žádná governance ani bezpečnostní pravidla, vysoké riziko úniku dat



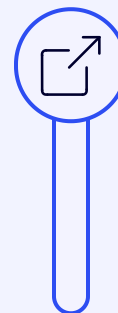
Úroveň 2: Řízené

Schválený seznam AI nástrojů, základní školení a bezpečnostní guidelines, sledování prvních případů použití



Úroveň 3: Integrované

AI začleněna do klíčových procesů, definované KPI a měření ROI, dedikovaný tým pro správu AI



Úroveň 4: Optimalizované

AI-first přístup ve všech procesech, vlastní přizpůsobené AI modely, automatická optimalizace a pokročilá analytika



Úroveň 5: Transformativní

AI jako základ business modelu, vlastní AI produkty a služby, vedoucí pozice na trhu v AI inovacích

Využití AI v auditu (příklad)

Fáze 1: Plánování auditu

Klíčové riziko: "Over-reliance" na AI doporučení → chybný audit scope

1	2	3
Analýza rizik Příklad: AI vyhodnotí historická data a navrhne rizikové oblasti Přínos: Cílené zaměření auditu na kritické oblasti Riziko: Přehlédnutí nových nebo emergentních rizik	Příprava checklistů Příklad: Automatické vytvoření kontrolních otázek dle ISA Přínos: Úspora času, konzistence napříč projekty Riziko: Generické otázky neodpovídající realitě klienta	Benchmark analýza Příklad: Porovnání finančních a nefinančních ukazatelů klienta s odvětvovými průměry a konkurencí. Přínos: Odhalení potenciálních problémů a příležitostí, které by jinak nebyly zjištěny. Riziko: Závislost na kvalitě a dostupnosti srovnávacích dat, riziko zobecnění bez ohledu na specifika klienta.

Fáze 2: Realizace auditu

Klíčové riziko: Nekritické přijetí AI výstupů → audit failure

1	2	3
Kontrola dokumentů Příklad: Analýza smluv, faktur, board minutes Přínos: 100% pokrytí místo vzorkování Riziko: Halucinace – AI "vidí" co tam není	Analýza dat Příklad: Detekce anomálií v transakcích Přínos: Rychlá identifikace odchylek Riziko: Falešné pozitivy (80%+ false alarms)	Testování kontrol Příklad: Analýza účinnosti interních kontrol, simulace scénářů selhání. Přínos: Posílení spolehlivosti interního kontrolního systému, snížení manuální práce. Riziko: Složitost nastavení AI pro složité kontrolní rámce, "black box" problém.

Fáze 3: Závěr a reportování

Klíčové riziko: Faktické chyby v reportech → ztráta důvěryhodnosti, právní důsledky

1	2	3
Sumarizace zjištění Příklad: AI shrne klíčové body z audit evidence Přínos: Konzistence, profesionální struktura Riziko: Automation bias – auditor přestane kriticky myslet	Tvorba zpráv Příklad: Návrh textu audit reportu Přínos: Rychlost, jednotný profesionální jazyk Riziko: Falešné citace nebo nesprávná čísla	Quality review Příklad: AI kontrola kvality auditní dokumentace, identifikace nedostatků. Přínos: Zvýšení kvality auditu, odhalení chyb před finalizací. Riziko: AI nemusí pochopit kontext specifického klienta, riziko falešných doporučení.

Klíčové doporučení: Využití AI v auditu by mělo být vždy v kombinaci s lidským úsudkem a kritickým myšlením, aby se maximalizovaly přínosy a minimalizovala rizika.

Podpora napříč celým auditním procesem

1

Quality review

Konkrétní příklad: Kontrola souladu s auditorskými standardy

Přínos: Snížení chyb a zvýšení kvality

Hlavní riziko: Přehlédnutí subtilních problémů

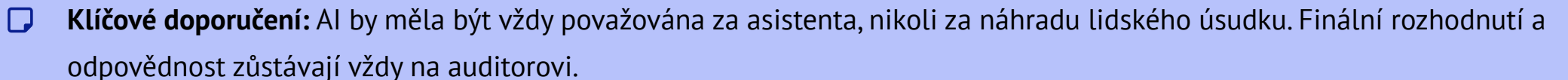
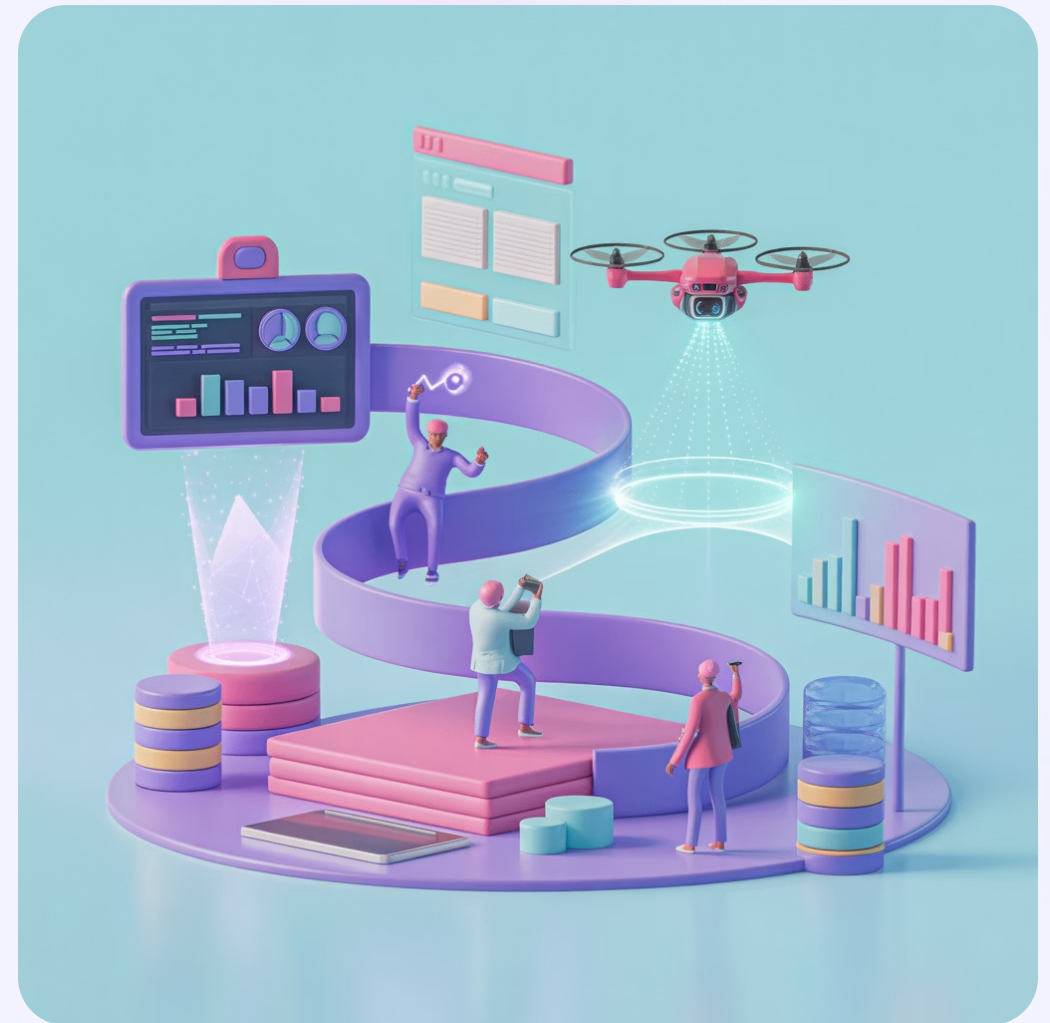
2

Knowledge management

Konkrétní příklad: AI asistent pro rychlé vyhledávání precedentů, technických memo, best practices z databáze firmy

Přínos: Demokratinizace expertního know-how,
úspora času při řešení netypických případů

Hlavní riziko: Tým aplikuje řešení z jiného jurisdikce nebo outdated guidance bez konzultace experta



Když výstupy AI „zní dobře“, ale mýlí se

Prakický test & důkaz existence halucinace: „Vypiš vyjmenovaná slova po N“

ChatGPT

Odpověď: Přidal nylon, nábytek,
pobýt, bydlet

Claude

Odpověď: Uvedl něha, něžný, změna,
měnit, peníze

Gemini

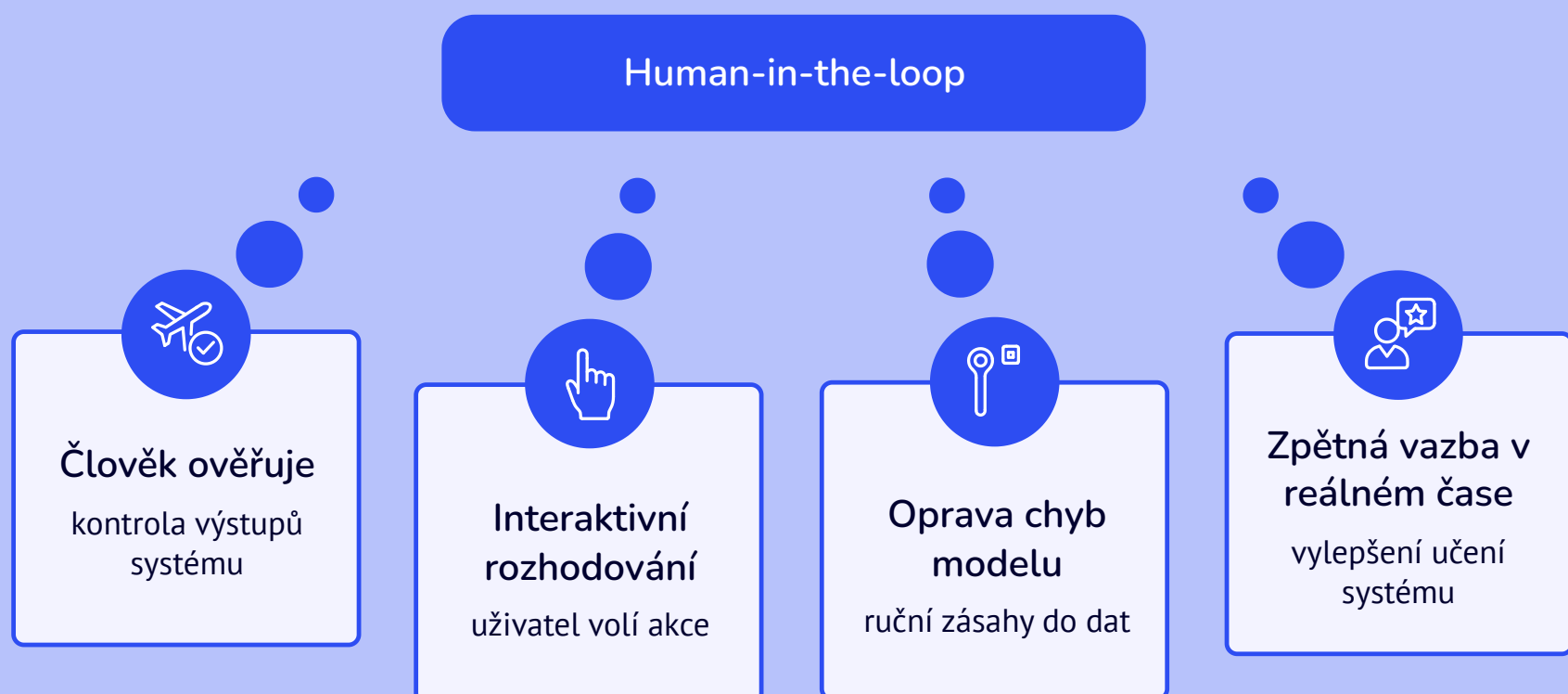
Odpověď: Vymyslel nyjší, nezývat,
někdejší

Co to znamená pro audit?

- AI může vygenerovat perfektně znějící zprávu s **neexistujícími čísly nebo fakty**
- Standardní "sense check" nestačí – potřebujete **source verification**

📄 **Poučení:** Bez lidské validace by „chytrá chyba“ mohla projít i skvěle provedeným auditem. Halucinace AI jsou sofistikované a přesvědčivé.

📄 Jak zajišťuje požadavek 'Human-in-the-loop'?



Základní kritéria pro výběr AI nástrojů

Jak vybíráte AI řešení? Co na to DPO? Používáte jednotný USE case formulář?

1

Ochrana klientských dat

End-to-end šifrování, izolace dat, mezinárodní certifikace (ISO 27001, SOC 2)

2

Soulad s regulací

Compliance s GDPR, připravenost na AI Act, soulad s ISA/ISQM standardy

3

Transparentnost a vysvětlitelnost

Schopnost vysvětlit, jak AI dospěla k výstupům a doporučením

4

Jasná odpovědnost

Definovaná odpovědnost za chyby – kdo nese riziko při selhání AI?

5

Cena a TCO

Celkové náklady na vlastnictví včetně školení, integrace a údržby

6

Stabilita dodavatele

Reference, finanční zdraví, dlouhodobá strategie a podpora

7

Lokalizace (bonus)

Podpora českého jazyka, místní technická podpora a znalost prostředí

Tři základní varianty provozování AI řešení

1

Public cloud

Výhody: Nízké náklady, rychlé nasazení, žádná údržba

Nevýhody: Riziko úniku dat, neauditovatelnost procesů

Příklady: ChatGPT, Gemini

Klientská data?: ❌ Ne

2

Privátní cloud

Výhody: Izolace tenantu, SLA, kontrola přístupu, audit trail

Nevýhody: Vyšší cena, nutnost governance

Příklady: Azure OpenAI, Google Vertex AI

Klientská data?: ⚠️ Podmíněně

3

On-premise řešení

Výhody: Plná kontrola, suverenita dat, nezávislost

Nevýhody: Vysoké nároky na správu a infrastrukturu

Příklady: Llama, Mistral

Klientská data?: ✅ Ano



⚠️ Kritické upozornění pro auditory:

- **AI jako služba:** Pokud vložíte klientská data do ChatGPT/Gemini, porušujete profesní tajemství
- **Privátní cloud:** Vyžaduje Business Associate Agreement (GDPR/NIS2 compliance)
- **On-premise:** Nejbezpečnější, ale náročné na IT resources a průběžné aktualizace modelů

Doporučení: Začněte s privátním cloudem kombinovaným se strict governance frameworkem a postupně vyhodnocujte možnost on-premise řešení.

Use Case – NotebookLM

Co to je?

- Asistent pro práci s vlastní interní dokumentací
- Generuje odpovědi s přesnými citacemi a zdroji
- Vytváří přehledy i audio shrnutí dokumentů
- Data zůstávají v rámci Google ekosystému

Technická omezení

Maximum 50 dokumentů najednou, není plně on-premise řešení, závislost na Google infrastruktuře

Kdy NEPOUŽÍVAT ✗

Klientská data bez šifrování

Riziko porušení profesního tajemství a důvěrnosti

Potřeba audit trail

Není možné sledovat, kdo co četl a kdy

Data s PII

Osobní identifikovatelné informace vyžadují vyšší ochranu



Use Case – DeepResearch

Co to je?

DeepResearch je pokročilý nástroj, který kombinuje analýzu velkých objemů dat s automatickým ověřováním zdrojů. Pomáhá při due diligence procesech a komplexních rešerších.

Hlavní přínosy

- Zvyšuje rychlost analytických procesů až o 70%
- Zajišťuje konzistenci výstupů napříč projekty
- Automaticky verifikuje a cituje zdroje informací
- Identifikuje vzorce a souvislosti v datech

Kdy NEPOUŽÍVAT ✗

→ Time-sensitive úlohy

Analýza může trvat i několik hodin – nevhodné pro urgentní situace

→ Požadavek 100% přesnosti

Stále existuje riziko halucinací a nepřesností v komplexních analýzách

→ Vysoce specializované otázky

Technické otázky mimo training data mohou vést k nesprávným závěrům

Sedm klíčových rizik využití AI v auditu



1. Halucinace

AI vygeneruje **falešná** čísla, neexistující citace, závěry nebo smyšlené "fakta"



2. Automation bias

Auditor **přestane kriticky myslet** a slepě důvěřuje AI výstupům



3. Porušení důvěrnosti

Únik citlivých klientských dat do nechráněných systémů



4. Nekvalitní vstupní data

Garbage in, garbage out – nekvalitní data vedou k chybným závěrům



5. Nevysvětlitelnost

Black box rozhodnutí bez možnosti auditovat logiku AI



6. Kybernetická bezpečnost

Útok na AI systém, manipulace s modely nebo data poisoning



7. Nesoulad s regulací

Rychlé změny v legislativě a regulaci (AI Act, ...)

AI incidenty a poučení z praxe

Rok	Firma	Popis incidentu	Dopad a následky
2025	Deloitte Australia	AI systém vygeneroval falešné citace a neexistující zdroje ve vládní auditní zprávě https://fortune.com/2025/10/07/deloitte-ai-australia-government-report-hallucinations-technology-290000-refund/	Kompletní revize zprávy, vrácení honoráře, reputační škoda
2025	FRC UK kontrola	Šest velkých auditních firem nemělo zavedeno sledování dopadu AI na kvalitu auditu	Regulační výtky, povinnost zavést monitoring do 6 měsíců



- Závěr:** Nízký počet zaznamenaných incidentů **≠ nízké riziko**. Absence systematického měření a reportování znamená, že existuje obrovská **slepá zóna** v pochopení skutečných rizik AI v rámci auditu.
Mnoho incidentů AI v auditu pravděpodobně zůstává neodhaleno nebo není veřejně reportováno kvůli obavám o reputaci a odpovědnost.

Sedm (vybraných) otázek pro auditory

1. Validací checklist

Máte povinný checklist pro validaci každého AI výstupu? Zahrnuje kontrolu citací, čísel a logické konzistence?

2. Definice odpovědnosti

Je jasné definováno, kdo nese odpovědnost za chybu AI? Partner, manager, nebo AI vendor? Kde končí odpovědnost technologie?

3. Hranice využití

Kdy AI pomáhá versus kdy nahrazuje auditora? AI například nesmí dělat "final judgement" bez lidského schválení.

4. Měření úspěšnosti

Měříte % auditních zjištění, kde AI selhala nebo naopak pomohla? Máte data pro informovaná rozhodnutí?

5. Komunikace s klienty

Jak informujete klienty o využití AI? Engagement letter, verbal disclosure, nebo písemný souhlas?

6. Detekce halucinací

Jak detekujete halucinace AI v reálném čase? Automatické red flags, peer review, nebo source verification?

7. Audit AI systémů

Máte auditora vyškoleného na auditování AI systémů?



Call to Action: Chcete být v obraze? Odebírejte Newsletter AI007 – https://lnkd.in/e8z_F9Rs

Sedm opatření v rámci AI governance



AI Governance Board

Pravidelné měsíční schůzky s reportováním přímo CEO nebo Audit Committee. Složení: partner, IT, risk manager, legal counsel.



100% human-in-the-loop pro klientské výstupy

AI navrhuje a připravuje, ale člověk vždy schvaluje. Žádný dokument určený klientovi nesmí být publikován bez lidské validace.



AI Tool Register

Centrální evidence: Nástroj | Vendor | Použití | Risk Rating | Schvalovatel | Datum schválení | Review date



DLP + monitoring

Data Loss Prevention – automatická blokáce citlivých dat (rodná čísla, čísla účtů, PII). Real-time monitoring AI dotazů.



Metriky kvality AI

Pozitivní metriky: % AI-assisted auditů | Čas ušetřený | Počet dokončených auditů

Rizikové metriky: False positive rate | Near-miss incidenty | AI halucinace v QC | Čas na korekce AI chyb | Aktivace AI pause button

Reporting: Měsíční dashboard + quarterly trend analýza. Red flags: >30% false positives nebo >3 near-misses/měsíc → eskalace do Governance Board.



Povinná školení každých 3-6 měsíců

Certifikované školení pro všechny auditory zahrnující technické aspekty, rizika a praktické case studies. Povinný "refresh training" při zavedení každého nového AI nástroje (nečekat na pravidelné školení).



"AI pause button"

Jasně definovaný proces, jak okamžitě zastavit použití AI při detekci problému. Eskalační postupy a odpovědnosti.



Doporučení: Kombinujte vhodná technická, procesní, organizační opatření

Doporučené informační zdroje



ISACA AI Audit Toolkit

Komplexní nástroje pro audit AI systémů
vydané v roce 2025

isaca.org/ai-audit-toolkit



MIT AI Risk Catalog

Taxonomie rizik umělé inteligence od
Massachusetts Institute of Technology

airisk.mit.edu



FRC Thematic Review

První tematický přehled o dopadu AI na
kvalitu auditu od UK regulátora

frc.org.uk/ai-guidance



IAASB Technology Publications

Staff publikace o technologii v auditu od
International Auditing and Assurance
Standards Board

iaasb.org/focus-areas/technology



NIST AI Risk Management Framework (AI RMF)

Framework pro řízení rizik umělé
inteligence od National Institute of
Standards and Technology

[nist.gov/itl/ai-risk-management-
framework](https://nist.gov/itl/ai-risk-management-framework)

Tyto zdroje poskytují aktuální guidance, případové studie a praktická doporučení pro bezpečné a efektivní využití AI v auditu.



Shrnutí

AI není budoucnost – je přítomnost

87% organizací už AI využívá. Otázka není zda, ale jak bezpečně a efektivně.

Začněte s AI governance, ne pouze s technologií

Pravidla a procesy musí předcházet nasazení nástrojů.

Bez governance představuje AI obrovské riziko.

Human-in-the-loop je základ

AI asistuje, člověk rozhoduje. Finální odpovědnost vždy zůstává na auditorovi.

Měřte, učte se, adaptujte - kombinujte rychlé a systematické kroky

90-denní roadmapa vám pomůže nastartovat AI governance ve firmě.

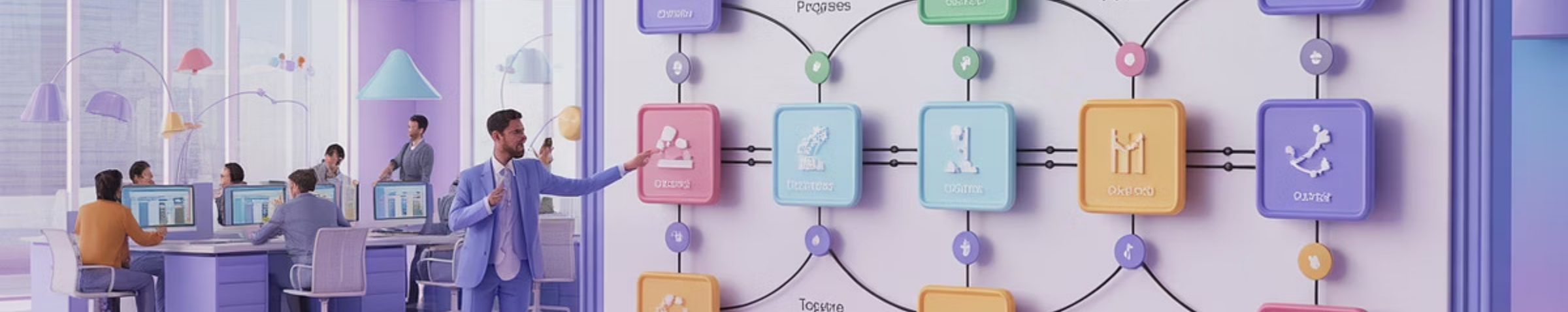
Vaše další kroky ?



Děkuji za pozornost! Otázky a diskuse u (kávy) jsou vítány.



Přílohy



90-denní AI Roadmap

Týden 1-2: Audit současného stavu

Klíčové aktivity:

- Anonymní průzkum týmu o současném využití AI nástrojů
- IT audit – detailní kontrola přístupů k AI službám a zabezpečení
- Identifikace top 3 rizik specifických pro vaši organizaci
- Mapování existujících AI nástrojů a jejich využití

Očekávaný výstup:

Komplexní report současného stavu obsahující:

- Přehled využívaných nástrojů
- Identifikovaná rizika
- Gap analýza proti best practices
- Doporučení pro další kroky

Týden 3-4: Quick wins

01

Schválit bezpečné nástroje

Identifikovat a formálně schválit 2-3 bezpečné AI nástroje (například NotebookLM, Microsoft Copilot s business licencí)

02

Vytvořit Quick Start Guide

Připravit jednoduchý průvodce na 1 stránce s praktickými příklady a bezpečnostními pravidly

03

První školení týmu

Zorganizovat 60minutový workshop zaměřený na praktické použití a základní bezpečnostní principy

☐ **Očekávaný výstup:** Tým má k dispozici legální a bezpečné nástroje s jasným návodem na použití. Eliminace shadow IT a neschválených nástrojů.

Měřitelné cíle:

- 100% týmu proškolen v základech
- Minimálně 50% týmu začne používat schválené nástroje
- Nulový počet bezpečnostních incidentů



Měsíc 2: Governance a struktura

AI Policy

Vypracovat stručnou AI Policy na maximálně 2 stránky A4 pokrývající:

- Povolené a zakázané nástroje
- Pravidla pro práci s klientskými daty
- Eskalační postupy při problémech
- Odpovědnosti a role

AI Champion

Jmenovat interního AI Championa s jasně definovanými odpovědnostmi:

- Koordinace AI iniciativ
- Pravidelné reportování vedení
- Podpora týmu v otázkách AI
- Monitoring trendů a aktualizací

AI Tool Register

Vytvořit a udržovat centrální registr AI nástrojů obsahující:

- Název nástroje a vendor
- Risk rating a klasifikace
- Schvalovatel a datum schválení
- Review date a podmínky použití

Očekávaný výstup: Jasná a vymahatelná pravidla pro celou organizaci s definovanými odpovědnostmi a procesy.

Měsíc 3: Pilotní projekt

Týden 1: Výběr projektu

Vybrat vhodný audit pro testování AI –
ideálně středně velký projekt s kvalitními
daty

1

2

3

4

Týden 3: Realizace

Provést audit s využitím AI, průběžně měřit
čas, kvalitu a náklady oproti baseline

Týden 2: Definice scope

Přesně definovat, kde a jak bude AI použita –
konkrétní use cases a očekávané výstupy

Týden 4: Evaluace

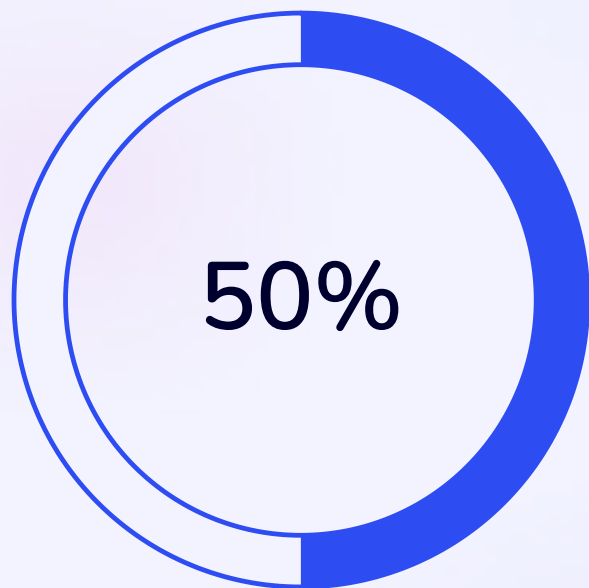
Detailně dokumentovat lessons learned,
problémy a úspěchy projektu

Klíčové metriky k měření:

- Čas strávený na jednotlivých fázích auditu
- Počet identifikovaných auditních zjištění
- Kvalita výstupů (peer review score)
- Spokojenost týmu a klienta
- Náklady na audit
- Počet AI chyb nebo halucinací
- Efektivita využití AI nástrojů
- ROI pilotního projektu

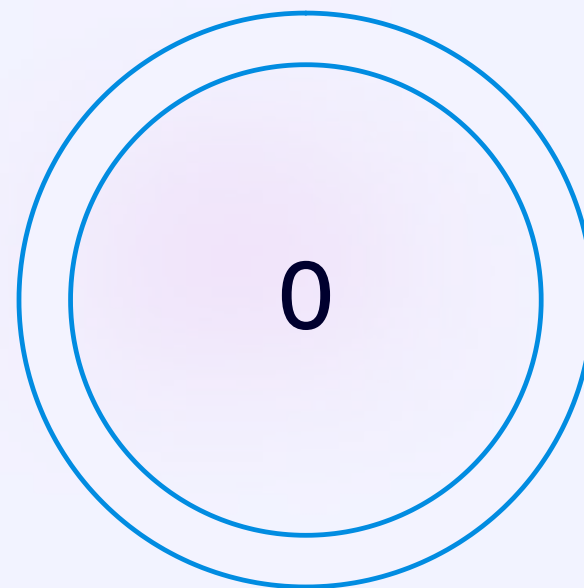
📄 **Očekávaný výstup:** Datově podložené rozhodnutí – **scale up** (rozšíření na další projekty), **pivot** (úprava přístupu), nebo **stop** (ukončení iniciativy).

KPIs po 90 dnech realizace



Adopce týmem

Minimálně 50% týmu aktivně používá schválené AI nástroje v každodenní praxi



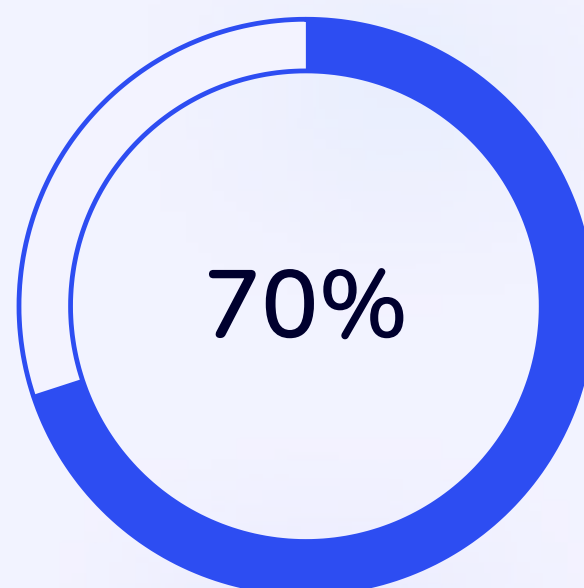
Bezpečnostní incidenty

Nulový počet bezpečnostních incidentů spojených s používáním AI



Ušetřené hodiny

Minimálně 100 hodin ušetřených díky efektivnějším auditním procesům



Spokojenost týmu

Průměrné hodnocení spokojenosti s AI nástroji minimálně 7 z 10 bodů

Další důležité metriky k vyhodnocení:

Kvalitativní ukazatele:

- Úroveň důvěry týmu v AI výstupy
- Počet eskalací nebo problémů
- Zpětná vazba od klientů
- Identifikované oblasti pro zlepšení

Kvantitativní ukazatele:

- ROI pilotních projektů
- Počet dokončených školení
- Využití jednotlivých AI nástrojů
- Čas strávený na routine tasks

Kritické rozhodnutí: Na základě těchto KPIs se rozhodněte o dalším směru AI strategie – pokračování, úprava nebo přehodnocení přístupu.