

ICT REVUE



Éra kvantových strojů

Proč kvantové počítače nejspíš nikdy nedosáhnou popularity umělé inteligence.

Riziko kódování

AI pomáhá vytvářet programovací kód rychleji. Často ale na úkor bezpečnosti dat.

Výkonnost AI

Firmy do umělé inteligence investují ve velkém. Jak efektivně ji ale využívají, měří český start-up.

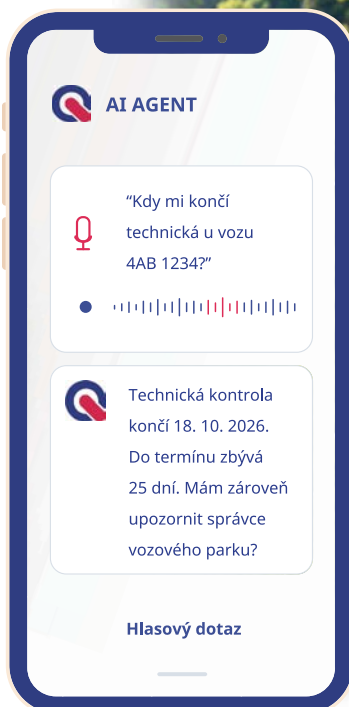
O ČEM MLUVÍTE SE SVOU FIRMOU VY?

Mluvte.

AI Agent poslouchá.

ERP má držet klíčová data, pravidla a firemní procesy. Sklad, doklady, portály, specializované aplikace a AI však někdy potřebují vlastní prostor.

V Quartexu navrhujeme, propojujeme a dlouhodobě rozvíjíme firemní systémy tak, aby pracovaly jako jeden funkční celek.



STAČÍ SE ZEPTAT HLASEM.

Položte otázku hlasem a získejte odpověď okamžitě.



DOKLADY
Q.Invoice



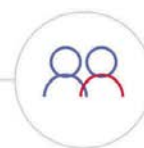
SKLAD
Q.WMS



REPORTY
Přehledné
informace



AI AGENT
Odpoví z dat,
která sám vidět



**APLIKACE
A PORTÁLY**
Zakázkové
řešení



ERP
Business Central
Dynamics NAV

BUSINESS CENTRAL

DYNAMICS NAV

Q.WMS

Q.INVOICE

APLIKACE

REPORTING

AI AGENTI

VAŠE FIRMA ZNÁ ODPOVĚĎ.

STAČÍ SE ZEPTAT.

www.quartexp Praha.cz

QUARTEX

JEDEN PARTNER PRO ERP A FIREMNÍ PROCESY.
OD ROKU 1998.

OBSAH

Kvantový počítač versus AI

04–06

O umělé inteligenci se dlouho jen mluvilo, aby pak během dvou let zcela změnila způsob každodenní práce. Kvantové výpočty vypadají jako podobný příběh. Platí pro ně ale zásadní odlišnost.



ERP, bezpečnost a cloud

09

Jaké požadavky vyplývají z nového zákona o kybernetické bezpečnosti pro provoz ERP systémů v cloudu a do jaké míry se v praxi naplňují?

AI v tvrdých datech

10–14

Firmy pumpují do AI spousty peněz, návratnost investice ale může být spekulativní. Jaká je realita, pomáhá měřit český start-up.



Riziko vibe codingu

16–19

Umělá inteligence dovede proces tvorby kódů velmi zrychlit. Velké nedostatky má ovšem v zabezpečení dat.

IT systém jako základ

20–21

Počítač lze nahradit snadno. Firma jako celek je ovšem na IT závislá víc než kdy dříve. Proto je důležité mít kvalitní návrh i správu systémů.



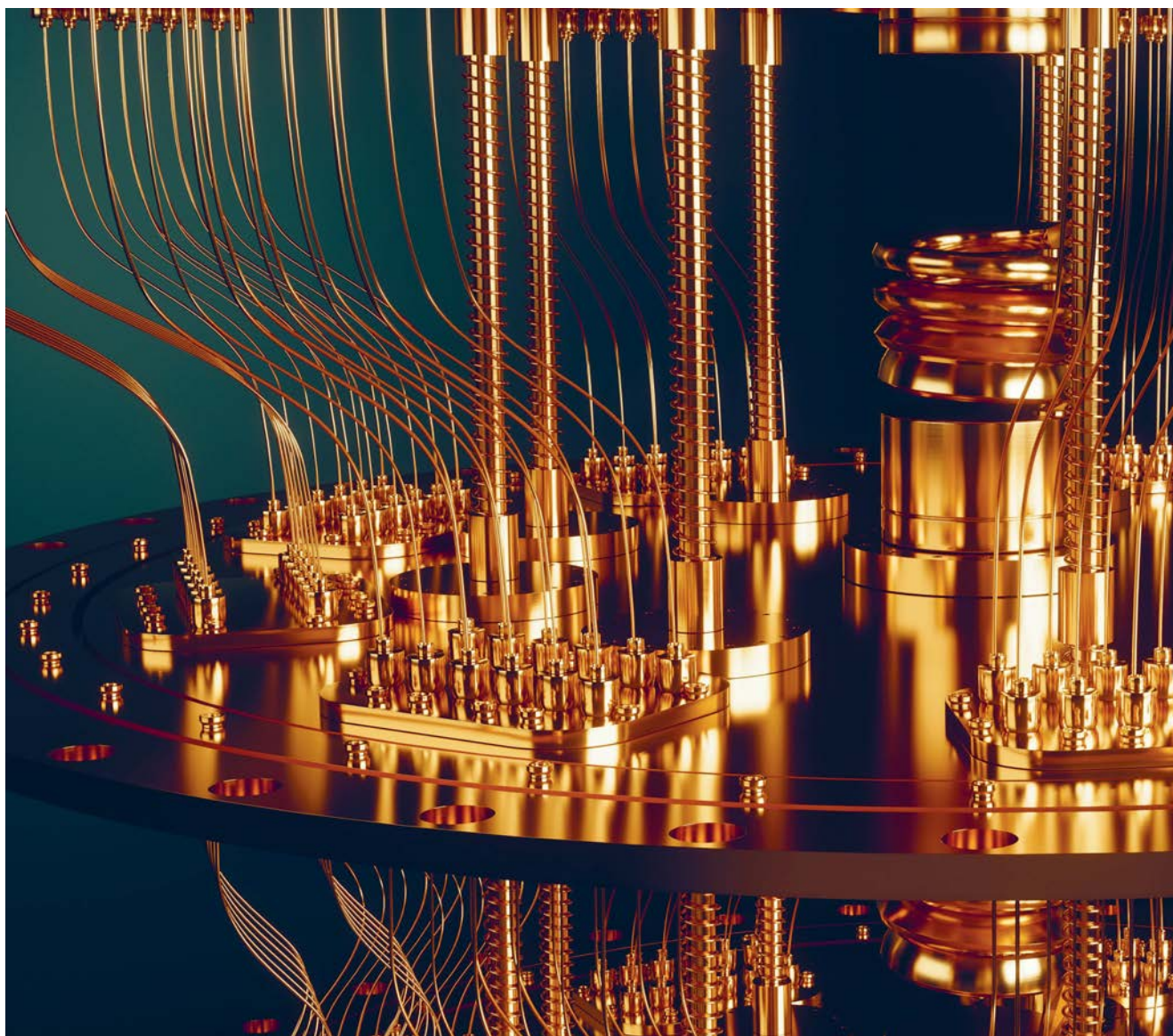
Strategie pro AI

22–26

Umělá inteligence je ve firmách téma natolik široké, že si zaslouží pozornost nejvyššího vedení. Pokud ji řeší jen IT, mohou vznikat projekty bez byznysové hodnoty, pokud jen byznys, trpí řízení a bezpečnost, říká v rozhovoru ředitel AI Transformation ve společnosti O2 Petr Hirš.

PROČ TO S KVANTOVÝMI POČÍTAČI NEBUDE JAKO S AI

O umělé inteligenci se dlouho jen mluvilo, aby pak během dvou let zcela změnila způsob každodenní práce. Kvantové výpočty vypadají jako podobný příběh, ale v jednom podstatném ohledu jsou přesným opakem. A právě z něj plyne rozdílný přístup firem.



Umělá inteligence byla desítky let tématem, které se nebralo příliš vážně. Dvakrát prošla obdobím nazývaným zima umělé inteligence: v polovině sedmdesátých let a znovu na přelomu osmdesátých a devadesátých let vyschlo financování poté, co obor nesplnil, co veřejně sliboval. Absolutní zlom pak nastal v roce 2022 spuštěním ChatGPT. Analogie s kvantovými počítači má reálný základ: algoritmus, jímž by kvantový počítač prolomil dnešní šifrování, publikoval americký matematik Peter Shor už v roce 1994. Jsme tedy dvaatřicet let na cestě ke stroji, který by to měl dokázat.

Podobnost ale končí právě u historie. Umělá inteligence se svezla na hardware, který už existoval. Konkrétně na herních grafických kartách. Byla univerzální, takže jeden model obsloužil ti-

Kvantový počítač nemá ambici nahradit běžné počítače. Doplní je spíše servery a superpočítače.

síce různých úloh. A šířila se zdola, protože si ji každý mohl vyzkoušet zdarma. Kvantový počítač nemá ani jednu z těchto vlastností. Je to účelové zařízení bez druhotného trhu, zvládá relativně úzký okruh úloh a spotřebitelské aplikace neexistují. Odpovídají tomu i čísla: porovnáme-li údaje Stanfordovy univerzity o investicích do umělé inteligence s daty poradenské společnosti McKinsey, přiteklo loni do kvantových firem zhruba třicetkrát méně soukromého kapitálu než do AI.

Podstatnější je ovšem odlišnost v samotném zadání. U umělé inteligence technologie fungovala a hledalo se, k čemu ji použít. U kvantových počítačů je to obráceně: víme přesně, k čemu se hodí, ale chybí stroj. Tím se vysvětluje i zdánlivý paradox celého tématu, totiž že jedno konkrétní kvantové riziko má pevné termíny, zatímco obchodní příležitost zatím nikoli.

Kvanta bez fyziky

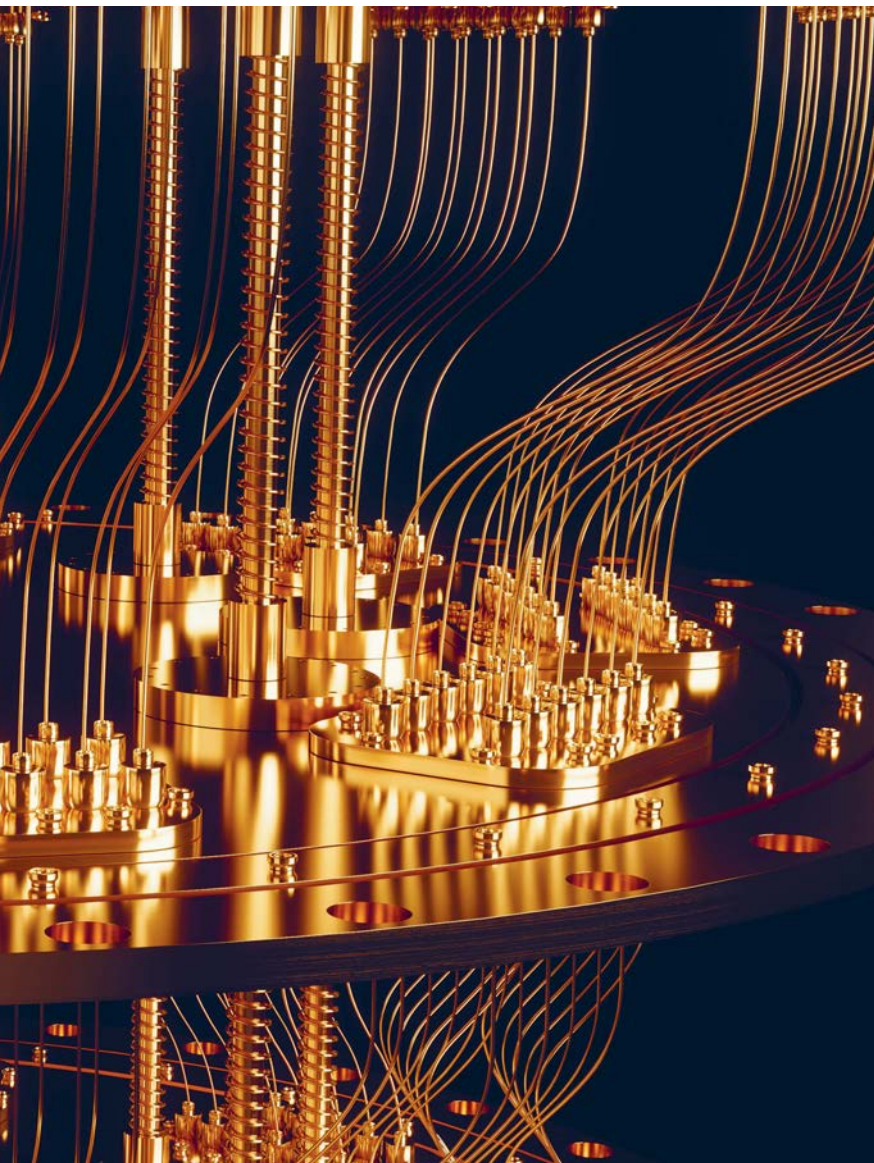
Kvantový počítač není rychlejší počítač. Ostatně v naprosté většině firemních výpočtů je beznadějně horší než běžný server. Jde o specializované zařízení, které doplňuje servery či superpočítače. Zatímco klasický systém odvede drtivou většinu práce, kvantovému procesoru předá jen malý, přesně specifikovaný podproblém se zvláštní matematickou strukturou.

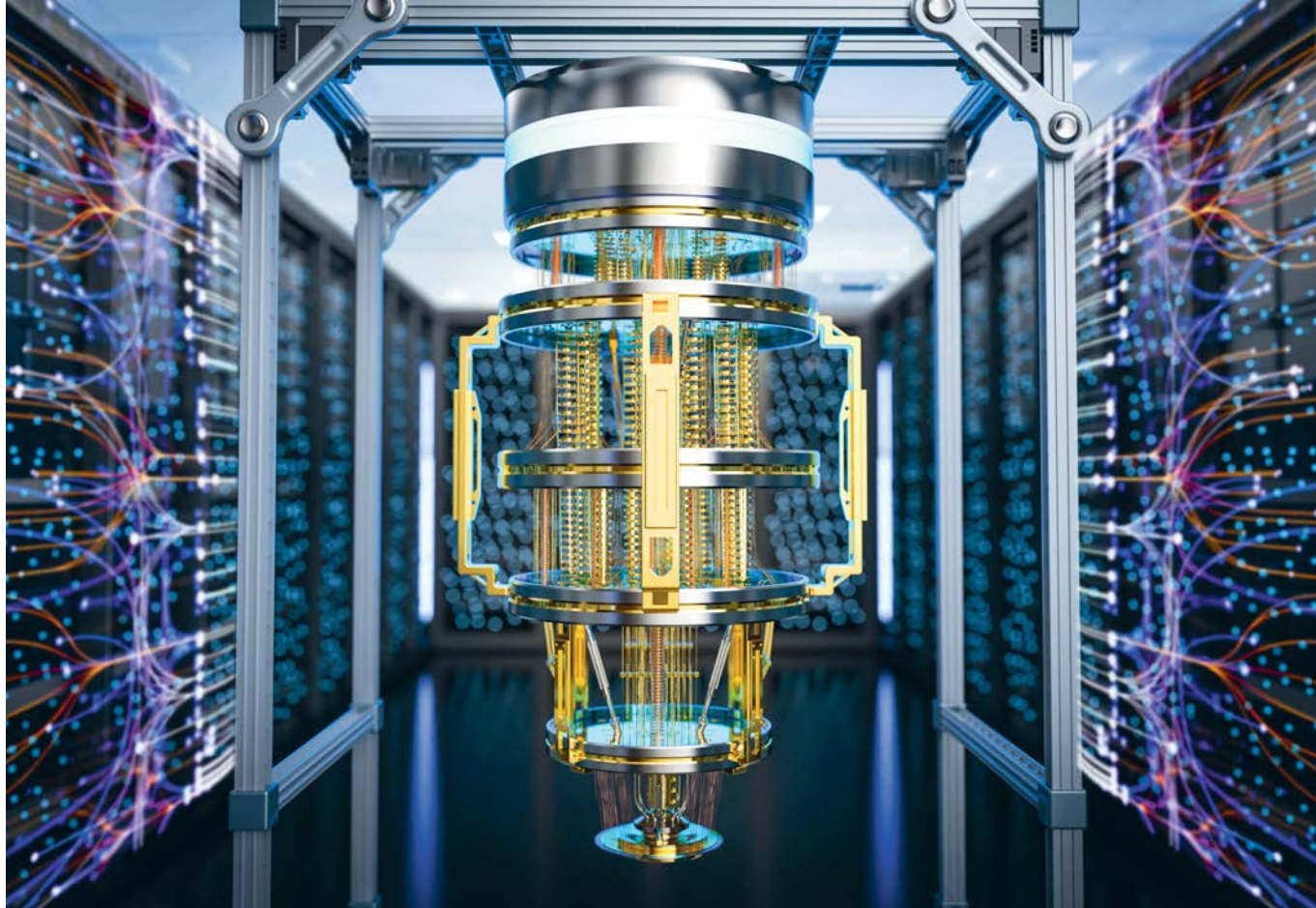
Údaj, kterým se v prezentacích mává nejčastěji, tedy počet qubitů, sám o sobě nevypovídá téměř o ničem. Qubit je základní jednotka takového stroje, obdoba bitu v běžném počítači, jen s výrazně vyššími nároky na prostředí a s podstatně vyšší chybivostí. A právě chybivost je, spolu s rozdílem mezi qubitem fyzickým a logickým, rozhodující. Aby stroj zvládl delší výpočet bez ztráty výsledku v šumu, spojují se dnes desítky až stovky fyzických qubitů do jednoho spolehlivého logického. Mezi strojem se stovkou fyzických a strojem se stovkou logických qubitů proto leží několik řádů investic a několik let vývoje. Srovnávat kvantové počítače podle počtu qubitů tedy nedává smysl.

Nekupuje se stroj, ale jeho čas

Kvantové počítače se do firemních datových center nepořizují a v dohledné době ani pořizovat nebudou. Nároky na chlazení a údržbu z nich dělají zařízení bližší urychlovači částic než serveru a provozovat je umí několik desítek pracovišť na světě, která poskytují jejich strojový čas prostřednictvím cloudu. Od letošního července je k dispozici i první takový stroj v Česku. Tento kvantový počítač provozuje národní superpočítačové centrum IT4Innovations v Ostravě. Vznikl v rámci EuroHPC, společného programu Evropské unie a členských států pro budování superpočítačové infrastruktury, otevřené i podnikům.

Zatímco strojový čas kvantových počítačů je relativně dostupný a jedná se o provozní náklad, problém je se specialisty, kteří dokážou obchodní úlohu přeložit do podoby srozumitelné kvantovému procesoru. A nákladná je také příprava vstupních dat.





Co vlastně umí a co ne?

Doménou kvantových počítačů jsou specifické úlohy, jako například simulace molekul a materiálů. Chování elektronů je kvantový jev a klasické výpočty na některé třídy sloučenin nestačí. Proto je kvantové počítání vhodné pro obory jako chemie, farmacie nebo výzkum materiálů. Cílem ovšem není získat konkrétní výsledek, ale zúžit počet kandidátů pro další laboratorní výzkum a testy, tedy zkrátit nejdražší fázi vývoje.

Marketingově oblíbenou oblastí je například optimalizace plánování tras nebo rozvrhování výroby. Poptávka je tu největší, protože takové úlohy řeší mnoho podniků, ale prokazatelné výsledky a praktické přínosy zatím chybí.

Nůžky mezi klasickými a kvantovými výpočty se ale příliš rychle nerozevírají. Při srovnávání výkonu hovoříme o takzvané kvantové výhodě, tedy o stavu, kdy kvantový stroj vyřeší konkrétní úlohu prokazatelně lépe než nejlepší známý klasický postup. Navíc nejde o stav, kterého bude dosaženo jednou provždy. Neustále se totiž zlepšují i klasické algoritmy a řadu efektních demonstrací minulých let posléze dostihly chytřejší postupy na běžném hardwaru. Je proto důležité se ptát, s jakou klasickou metodou byl výsledek porovnán, kdo ji naprogramoval a zda šlo o skutečná provozní data, nebo o testovací sadu.

Trh s kvantovými počítači ale bezpochyby existuje. Podle dubnové studie Quantum Technology Monitor 2026 poradenské společnosti McKinsey spolupracuje s dodavateli kvantových technologií přes tři sta organizací a jejich tržby

Kvantové počítače si umí poradit především se specifickými úlohami, jako je simulace molekul a materiálů.

loni poprvé přesáhly miliardu dolarů, s výhledem na 4,4 miliardy v roce 2028. Jde ovšem o tržby dodavatelů, nikoli o úspory zákazníků. Podle McKinsey by mohly kvantové technologie do roku 2035 vytvořit hodnotu 1,3 až 2,7 bilionu dolarů. Jde ovšem o modelový výpočet přínosů napříč odvětvími, nikoli o prognózu.

Riziko s jasnými termíny

Už více než tři desítky let víme, jakou hrozbu představuje kvantový počítač pro dnešní metody šifrování. Ohrožena je asymetrická kryptografie, tedy šifrování, na kterém dnes stojí důvěryhodnost internetových spojení, elektronických podpisů i plateb. Útočníci přitom na kvantový počítač čekat nemusí: stačí, aby už dnes ukládali odposlechnutou komunikaci a dešifrovali ji až ve chvíli, kdy bude k dispozici dostatečně výkonný kvantový stroj.

Riziko označované jako „harvest now, decrypt later“ se proto týká všech informací, které musí zůstat důvěrné ještě v příštím desetiletí. Může jít o smlouvy, dokumentaci k know-how, komunikaci a mnoho dalšího. Takzvaný Q-Day, tedy hypotetický den, kdy bude spuštěn dostatečně silný kvantový počítač schopný prolomit asymetrické šifry jako RSA nebo ECC, by podle každoročního průzkumu Global Risk Institute mohl nastat už po roce 2030.

Řešení přitom už existuje. Americký institut pro standardy NIST, jehož normy fakticky probírá celý trh, vydal už v srpnu 2024 první sadu postkvantových algoritmů, tedy šifer navržených

tak, aby je kvantový počítač nedokázal prolomit. Otázkou tedy není, jakou technologii šifrování nasadit, ale jak dlouho potrvá výměna současných algoritmů.

Že jde o reálné riziko, potvrzuje i americký exekutivní příkaz z letošního června, který ukládá federálním úřadům převést nejcitlivější systémy na postkvantové šifrování do konce roku 2030 a digitální podpisy do konce roku 2031. Tyto požadavky mají přejít na dodavatele federální správy a jejich prostřednictvím i na celý dodavatelský řetězec, evropské firmy nevyjímaje.

Členské státy Evropské unie se navíc už loni shodly na společném postupu, podle kterého má být kritická infrastruktura převedena na postkvantové šifrování rovněž do konce roku 2030. Velcí provozovatelé internetové infrastruktury ani na regulaci nečekají: Google či společnost Cloudflare, přes kterou prochází podstatná část světového internetového provozu, oznámily dokončení vlastního přechodu do roku 2029.

Kde začít

Nejtěžší na výměně šifrování není nasazení nového algoritmu, ale zjištění, kde všude se ve firmě stávající metody šifrování používají. Kryptografie se do systémů integrovala nejméně tři desítky let, včetně knihoven zabudovaných v aplika-

“

**Už více než
tři desítky
let víme,
jakou hrozbu
představuje
kvantový
počítač pro
dnešní metody
šifrování.**

cích od dnes už neexistujících dodavatelů nebo certifikátů v zařízeních, která už dávno nikdo neaktualizuje. Proto je nutné začít důkladnou inventurou.

V případě velkých organizací se pak přechod na kvantově odolné šifrování odhaduje na deset i více let. Už dnes je také nutné řešit způsob šifrování při pořizování strojů a zařízení s životností deset a více let, ať už jde o výrobní technologie, měřicí a zdravotnická zařízení nebo třeba platební terminály.

Vzhledem k dostupnosti prvních kvantových počítačů začíná být aktuální také otázka jejich využití v podnikové praxi. Rozhodující je, zda existuje konkrétní úloha, která vyžaduje kvantové výpočty, protože na ni tradiční stroje nestačí. Současně by ale měla být natolik důležitá, aby se využití kvantového počítače pozitivně promítlo do finančních výsledků. V každém případě se vyplatí obor kvantových výpočtů pečlivě sledovat.

V tomto ohledu je užitečná zkušenost s umělou inteligencí. Nejvíce totiž z jejího nástupu vytěžily firmy, které měly pořádek v datech, nikoli ty, které nejrychleji nakoupily modely. Ekvivalentem uklizených dat je u kvantových výpočtů přehled o vlastní kryptografii a schopnost vyměnit ji bez přepisování aplikací.

Inzerce

minerva.

**ERP je základ ekosystému vaší firmy.
ERP QAD Adaptive pro průmyslová odvětví s agentní AI
dodá vašemu podnikání nový rozměr.**



**Váš úspěch začíná
u správného systému.**



www.minerva-is.eu



marketing@minerva-is.cz



**Kontaktujte nás
pro více informací a začněte
transformovat svou výrobu.**



Bezpečná 5G mobilita i v prostředí s nebezpečím výbuchu

Digitalizace průmyslu se netýká jen výrobních linek a strojů. Stále více mění také práci lidí přímo v provozu, jehož automatickou součástí bývají mobilní zařízení. V prostředí s nebezpečím výbuchu však běžný smartphone použit nelze. Mobilní technologie zde musí splňovat specifické bezpečnostní požadavky a zároveň nabídnout moderní konektivitu včetně 5G.

Společnost Pepperl+Fuchs, kterou průmysl zná již více než 80 let hlavně jako výrobce technologií pro automatizaci a ochranu prostředí s nebezpečím výbuchu, rozšířila své portfolio také o mobilní zařízení určená právě pro náročné průmyslové a Ex provozy (pracovní oblasti a provozy s nebezpečím výbuchu plynů, par, prachů nebo vláken – pozn. red.). Jejich nová generace 5G smartphonů a tabletů propojuje nekompromisní bezpečnost s flexibilitou moderní mobilní kanceláře a otevírá nové možnosti pro průmyslové IT.

Nebezpečí ukryté v kapse

V moderních průmyslových provozech se mohou vyskytovat hořlavé plyny, páry nebo prach, které ve směsi se vzduchem vytvářejí potenciální riziko vzniku výbuchu. V těchto Ex zónách proto platí přísná pravidla pro používání elektrických a elektronických zařízení. Běžný mobilní telefon pro takové prostředí není konstruován ani certifikován. „Je zde nutné používat zařízení s odpovídající Ex certifikací,“ vysvětluje Martin Koráb, zástupce společnosti Pepperl+Fuchs.

„Běžná elektronika totiž není konstruována tak, aby bylo za každé situace vyloučeno riziko vzniku jiskry či vysoké povrchové teploty, a ty by mohly ve výbušném prostředí vyvolat explozi,“ říká Koráb s tím, že do Ex zón proto mohou pouze zařízení odpovídající požadavkům pro konkrétní prostředí a způsob použití. „U mobilních zařízení se vedle samotné Ex ochrany řeší také mechanická odolnost, antistatické vlastnosti a v současnosti stále více i kybernetická bezpečnost,“ dodává Koráb.

Mobilní kancelář ve výbušném prostředí

Potřeba přenosu dat a přístupu k podnikovým systémům, jako jsou MES či ERP, roste i v rizikových provozech. I sem proniká koncept mobilních kanceláří, kdy mobily a tablety přistupují ke stejným informačním systémům jako kancelářské počítače, a proto se stávají plnohodnotnou součástí podnikové IT infrastruktury. Současně roste význam privátních 5G sítí, které umožňují rychlý a spolehlivý přenos dat i v rozsáhlých výrobních areálech a vytvářejí základ pro rozvoj průmyslu 4.0 a IoT.



Portfolio mobilních zařízení Pepperl+Fuchs reaguje na tyto požadavky kombinací odolného hardwaru a digitálních funkcí. Smartphone Smart-Ex® 03 nabízí biometrické ověřování pomocí otisku prstu a režim Desktop Mode, který po připojení k příslušnému hardwaru umožňuje využívat telefon jako pracovní stanici. Uživatel pracuje se stejnými aplikacemi a daty v provozu i v kanceláři bez nutnosti jejich přenosu mezi více zařízeními.

Nově přidává podporu sítí 5G, má intuitivní ovládání, disponuje funkcí „Glove Mode“ pro bezproblémové ovládání v rukavicích či „Smart Call Handling“ pro správu hovorů pomocí gest. Nabízí i možnost rozšíření o čtečku čárových kódů. Je vhodný pro údržbu, logistiku i správu majetku.

Tablet Tab-Ex® 05 s větším displejem se stává branou do světa rozšířené reality díky certifikaci Google ARCore a integrovaným senzorům. Lze jej využít při inspekcích, vzdálené podpoře, školení pracovníků nebo prediktivní údržbě.

S rostoucí digitalizací průmyslu se mobilní zařízení mění z komunikační pomůcky na důležitý pracovní nástroj. „Budoucnost nebude stát pouze na vyšším výkonu zařízení, ale především na jejich schopnosti bezpečně propojit zaměstnance, technologie a podnikové informační systémy. Právě spojení ochrany proti výbuchu, 5G komunikace, bezpečné správy zařízení a digitálních trendů představuje jeden z hlavních trendů moderní průmyslové automatizace,“ shrnuje Martin Koráb.

PEPPERL+FUCHS

■ **Globální technologická společnost se sídlem v německém Mannheimu patří mezi přední světové výrobce průmyslových senzorů a technologií pro automatizaci a ochranu prostředí s nebezpečím výbuchu.**

■ **Společnost, založená v roce 1945, se zaměřuje na factory automation a process automation a její řešení nacházejí uplatnění například ve strojírenství, automobilovém průmyslu, intralogistice, procesním průmyslu nebo mobilních aplikacích.**

■ **Vedle senzorové techniky se Pepperl+Fuchs stále více zaměřuje také na digitalizaci, průmyslovou komunikaci a propojení fyzických procesů s digitálními systémy.**

Jaké požadavky vyplývají z nového zákona o kybernetické bezpečnosti pro provoz ERP systémů v cloudu a do jaké míry se v praxi naplňují?



Martin Dudek
solution architect,
IT Specialist, SAP

Nový zákon o kybernetické bezpečnosti přináší na firmy nový regulační tlak, který se promítá mimo jiné do podmínek smluv SLA (Service Level Agreement – pozn. red.). SAP k implementaci požadavků zákona přistupuje na dvou úrovních: platformové a aplikační.

Na platformové úrovni jsou bezpečnostní požadavky v cloudových řešeních SAP součástí systému už od jeho návrhu. SAP využívá takzvaný přístup „defense in depth“, tedy vícevrstvou ochranu, která zahrnuje zero-trust architekturu, řízení identit a přístupů, správu zranitelností, bezpečnostní logování a monitoring i business continuity a disaster recovery.

SAP každoročně absolvuje přibližně 130 interních a externích auditů a udržuje klíčové certifikace, které podporují soulad s evropskou směrnicí o kybernetické bezpečnosti (NIS2) a zákonem o kybernetické bezpečnosti (ZoKB).

Na aplikační úrovni SAP nabízí specializované bezpečnostní nástroje. Zákazník pak zodpovídá především za nastavení bezpečnostních rolí, správu identit firemních uživatelů a integraci se svým vlastním SIEM nástrojem. Právě tato oblast bývá v praxi největší výzvou. SAP proto zákazníky aktivně podporuje prostřednictvím svých nástrojů a metodiky, která zahrnuje i samostatné fáze zaměřené na bezpečnostní plánování.

Více odpovědnosti pro provozovatele

U cloudového řešení zákazník přenáší velkou část technické odpovědnosti za bezpečnost infrastruktury na SAP jako poskytovatele. Zákon v takovém případě vyžaduje jasný model sdílené odpovědnosti, který SAP transparentně dokumentuje smluvně i technicky. U on-premise řešení leží celá tíha plnění zákonných povinností na straně provozovatele systému.

V souladu se zkušenostmi z trhu platí, že přechod zákazníků do cloudu a zavedení strukturovaného monitoringu přináší lepší přehled o bezpečnostních incidentech. Vyšší počet hlášených incidentů tak často odráží především lepší schopnost jejich detekce, nikoli zhoršení celkové bezpečnostní situace.

Povinnost hlásit incidenty Národnímu úřadu pro kybernetickou a informační bezpečnost (NÚKIB), kterou ukládá zákon, tuto transparentnost dále posiluje. Cloudové nástroje SAP jsou navrženy tak, aby zákazníkům generování potřebných reportů usnadnily.



Vladimír Karpecki
senior konzultant,
Minerva Česká republika

Z pohledu dodavatele ERP je potřeba rozlišit dvě roviny – samotný software a jeho provoz. Minerva implementuje ERP QAD, jehož tvůrcem je americká společnost QAD. Za splnění požadavků zákona o kybernetické bezpečnosti na straně samotného softwaru proto odpovídá jeho autor. Stejně jako u ostatních světových ERP systémů i nejnovější verze QAD reflektují aktuální požadavky na kybernetickou bezpečnost.

Úkolem Minervy je pak implementovat tento ERP. U on-premise zákazníků dále pomoci zajišťovat požadavky zákona o kybernetické bezpečnosti a u zákazníků cloudových služeb Minervy trvale pokrývat požadavky na kybernetickou bezpečnost ve spolupráci se zákazníkem.

Realita je tedy taková, že se technické požadavky neřeší izolovaně pouze při jednorázové implementaci, ale promítají se do průběžné spolupráce mezi výrobcem ERP, implementačním partnerem, poskytovatelem cloudových služeb a samotným zákazníkem.

Jiné dělení odpovědnosti

Právě v rozdělení odpovědnosti je i hlavní rozdíl mezi cloudem a on-premise. Primárně v tom, že v cloudu je většina zodpovědnosti na straně poskytovatele cloudových služeb (pořád je tady ale sdílená odpovědnost), zatímco v případě instalace on-premise je zodpovědný za kybernetickou bezpečnost zákazník, a dodavatel ERP mu poskytuje podporu. Požadavky zákona se tak v praxi liší především podle toho, kde končí odpovědnost dodavatele nebo poskytovatele služby a kde začíná odpovědnost zákazníka, zejména u provozních procesů, přístupových práv, interních pravidel a návazných podnikových systémů.

Realita v praxi zatím ukazuje, že samotný ERP nebývá nejčastějším zdrojem problémů. Počet registrovaných incidentů se podle informací, které máme, zvýšil, ale většinou jsou to kybernetické incidenty spojené s jinými podnikovými systémy.

Abyste situaci vztahující se ke splnění zákonných opatření, která se týkají kybernetické bezpečnosti, byla ještě složitější, některé subjekty jako třeba poskytovatelé cloud computingu mohou přímo podléhat prováděcímu nařízení EK 2024/2690. To plně nahrazuje technické požadavky české vyhlášky o bezpečnostních opatřeních (vydávané NÚKIB). V praxi to znamená, že při posuzování požadavků versus reality u ERP v cloudu nestačí sledovat pouze českou legislativu, ale je nutné zohlednit i evropský regulační rámec a konkrétní roli jednotlivých subjektů v dodavatelském řetězci.



**AI FIRMĚ POMÁHÁ?
REALITA MŮŽE
PŘEKVAPIT**

Firmy pumpují obrovské sumy do umělé inteligence, návratnost takové investice ale může být na míle vzdálená od domněnek managementu. Běžně se produktivita při vývoji softwarů počítá ukazateli, jako je počet operací nebo řádků v kódu. Český start-up Navigara měří produktivitu inženýrských týmů a přínos AI nástrojů pomocí pokročilých AI agentů a analýzy dat. Tvrdými daty prokazuje návratnost investic do umělé inteligence. „Management reportuje, jak je situace skvělá a všichni AI používají. A potom přijdeme my a zjistíme, že zaměstnanci sice AI používají, ale třeba jen pět minut denně, a ještě na něco, co vůbec nedává pro firmu smysl,“ říká **Jiří Bachel**, zakladatel start-upu Navigara.

B

Budou ještě potřeba ajťáci?

Budou. Z toho, co dnes vidíme, AI pomáhá vývojářům být lepší. A i když firmy zapojují umělou inteligenci primárně proto, aby snížily náklady na vývoj softwaru, zatím není v takovém stavu, aby se dala zautomatizovat veškerá práce. Díky AI zvládnou vývojáři víc práce. Někteří o trochu, někteří pětinasobně, někteří stokrát víc. Rozdíly mezi nimi jsou obrovské, ale stále budou potřeba.

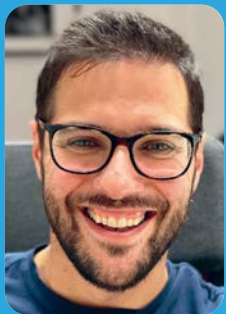
Jak moc je softwarový vývoj nečitelná položka pro vedení?

Ve firmách mu moc lidí nerozumí. Většinou se po vývojářích něco chce, například aby vyrobili nějakou funkci nebo dodali určitou vlastnost do aplikace, ale málokdo ví, jak je to náročné. Je slo-

žitě říct, že jsme rychlejší, když nepoznáte rozdíl a vidíte jen, že jste chtěli určitou věc a její vývoj trval měsíc. Těžko pak určíte, o kolik je firma rychlejší díky tomu, že utratila stovky tisíc dolarů za AI nástroje.

Když tedy firma zavede AI jako pomocníka v IT, má tendenci zveličovat její dopad na produktivitu zaměstnanců?

Tohle je zajímavé, protože ve chvíli, kdy firma zavádí AI nástroje, je pro její manažery přirozené říkat, jak AI pomáhá a vše je skvělé. Říkají to všichni. Liší se jen v tom, jak to podkládají daty. V momentě, kdy se ptají zaměstnanců, děje se to samé. Pro zaměstnance je těžké říct, že AI nepoužívají a je za ně zbytečná. Abyste se něco takového odvážili říct, musíte mít ve firmě silné postavení.



Jiří Bachel (39)

Společnost Navigara založil v roce 2022. Jeho software využívá například internetový vyhledávač letenek Kiwi nebo přední výrobce bezpečnostního softwaru Eset. V minulosti se podílel na založení dalších start-upů jako Onetone.AI nebo LOLO.ZONE. Dříve pracoval jako systémový inženýr, vystudoval ČVUT v Praze.



“

Někdy zjistíme, že zaměstnanci sice umělou inteligenci používají, ale třeba jen pět minut denně.

Výsledkem je, že se top managementu reportuje, jak je situace skvělá a všichni AI používají. A potom přijdeme my, podíváme se na data a často nevidíme rozdíl. Jsou momenty, kdy pochybujeme, jestli měříme správně.

Čím to bylo?

Nakonec zjistíme, že zaměstnanci sice umělou inteligenci používají, ale třeba jen pět minut denně, a ještě na něco, co vůbec nedává pro firmu smysl. Jakmile se začneme dívat detailněji na každý tým a na každého vývojáře, vidíme rozdílné stavy AI transformace. Na druhou stranu si ale firmy neuvědomují, jak náročný proces to je, když najednou mění styl práce stovkám lidí. Na výběr máte ze spousty úrovní a musíte se naučit mnoho technik. Nejen v kódování, ale i v tom, jak kód kontrolujete a vytváříte produktové zadání. A právě tady vidíme obrovské rozdíly mezi firmami. Některé zvládly trochu rychlejší kódování, ale začaly mít problém s kontrolou. Ty,

kteří zvládly kontrolu, mají problém vytvářet zadání, aby zvládaly tempo vývoje. Protože když se vám 200 vývojářů zrychlí dvakrát, nezvládá to produkt. Jakmile se jedna část zrychlí, další začne brzdit ostatní. A i ona se musí zlepšovat a to zlepšování firmu bolí.

Vývojáři se tedy musí naučit s AI pracovat. Kdy má smysl měřit využívání AI a návratnost investice?

Záleží, jak rychle tím chcete projít. Čím dříve začnete věci měřit, tím jednodušší je se zlepšovat. Takže z mého pohledu co nejdříve.

Práce ajťáků je hodně subjektivní a obtížně se posuzuje. Jak bylo složité do tohoto subjektivního světa hodit tvrdá data?

Inženýring má pocit, že je velmi subjektivní něco tvořit. My ale víme, že naše měření je konzistentní napříč různými typy týmů a rozdílnými druhy práce. Proto data čteme dvěma způsoby. Jeden



Každému ukazujeme, kde se aktuálně nachází, kde jsou mezery a jak se může zlepšit, říká k měření dat Jiří Bachel.

z nich porovnává týmy a firmy mezi sebou, druhý srovnává sebe sama – konkrétní tým oproti tomu, jak fungoval předtím, než začal používat AI nástroje. Měřit to je velmi komplikované. Strávili jsme víc než rok a půl tím, že jsme vymýšleli, jak to udělat. A zjišťujeme, že nejnáročnější je změnit tu škálu. To znamená rozdílné jazyky, rozdílné typy týmů, rozdílné typy firem, rozdílné chování vývojářů – tak, aby to bylo konzistentní a použitelné napříč obrovským spektrem. Stále se učíme a vylepšujeme. S každým novým zákazníkem a s každým novým neobvyklým chováním.

Máte pocit, že roste efektivita AI agentů a jejich použití v praxi?

AI modely jsou čím dál lepší, ale rozhodně největší dopad na výkon týmů mají samotní vývojáři. V momentě, kdy se naučí používat umělou inteligenci správně, výrazně ho posunou. Pokud je vývojář průměrně dobrý, pomůže mu průměrně. Pokud je výjimečně dobrý, umí ho poslat úplně do nebe.

Můžete nám popsat, jak přesně funguje Navigara? Co vše kontroluje a čte ze zdrojového kódu?

Napřed se podíváme na zdrojový kód a na to, jak se firma a jednotlivé týmy měnily před a po zapojení AI. Každému pak ukazujeme, kde se aktuálně nachází, kde jsou mezery a jak se může zlepšit. Potom se podíváme, kolik to celé stálo. Projdeme licence Anthropicu, OpenAI a dalších dodavatelů a řekneme, kolik firmu stálo zrychlení o tolik a tolik procent. Následně se podíváme do roadmapy firmy a zjišťujeme, jestli práce, kterou vývojáři za ty peníze udělali, byla vůbec plánovaná. Velmi často se nám stává, že velká část práce v roadmapě vůbec není. Neznamená to nutně, že byla zbytečná, jen nikdo neřekl, proč bylo potřeba ji udělat. Tím se ukážou chyby v interních procesech a jak zrychlení inženýringu přináší nové problémy.

Inzerce

EK017223

Jeden systém pro celou firmu

Propojte výrobu, sklady, finance i další agendy

www.qi.cz



V podstatě je to takový AI audit.

Jen se to neustále mění. Problémy, které firmy měly v lednu, byly jiné než ty dubnové a jiné než ty, které mají nyní. AI transformace nikdy nekončí.

Jak daleko do historie můžete jít?

Ve vývoji se používá nástroj, který se jmenuje Git a v něm je veškerá historie projektu. Takže můžete jít klidně i dvacet let do minulosti. Jde jen o to, jestli to dává konkrétní firmě smysl, nebo ne.

Jak se mění dynamika softwarových týmů v momentě, kdy část vývoje přebírají autonomní AI agenti?

Hodně se to liší podle velikosti firmy. Budu teď mluvit hlavně o těch větších, které mají přes sto vývojářů. Zhruba pětina lidí je nadšená. Tyhle týmy už nástroje používaly, mají je rády, okamžitě je adoptují a vidíme u nich velké zrychlení. Asi 40 procent lidí je spíš laxních. Vyzkouší si to, ale úplně nevědí, jak s tím pracovat. Zbývajících 40 procent AI vůbec používat nechce. Nějakým způsobem bojují s managementem a dokazují mu, že to nefunguje. To je pro firmu lidsky nejtěžší část. Ve chvíli, kdy se podaří tuhle skupinu přesvědčit, otevírá se jí obrovský prostor pro růst. Zvláště pokud je edukace personalizovaná pro konkrétního člověka.

Zrychlují AI agenti skutečně práci týmů, nebo spíš generují technologický odpad?

AI je pořád jen nástroj. Pracujete-li s ní špatně, přidělá vám práci. Když dnes nechám AI napsat článek, napíše ho, ale bude nečitelný pro normální lidi. Podobné je to v kódování. Proto se vývojáři nejdřív učí s AI pracovat. Pokud vývojář AI agentovi věří až moc a nechá ho udělat úplně všechno, vytvoří spoustu kódu, kterému nikdo nerozumí. Začnou se objevovat chyby, které už ani umělá inteligence neumí opravit. Stojí to spoustu peněz a po měsíci až dvou se firma rozhodne projekt zahodit a začít úplně znova.

Za jak dlouho se firmám může vrátit investice do AI?

Já si především myslím, že pro firmy je aktuálně povinnost do AI investovat, protože neinvestovat představuje obrovské riziko. To, co my vidíme, je, že na začátku firmy netrápí peníze. Posílají libovolné množství peněz jenom na to, aby to vůbec někdo ve firmě začal používat. Částky ale najednou vylétnou do obřích rozměrů a vedení začne trápit, že nevidí návratnost.

Firmu řídíte z amerického San Franciska. Jak je to pro váš byznys klíčové?

Já to rád přirovnávám k tomu, jako když vám v Česku někdo něco prodává z Kazachstánu. Neříkám, že by nešlo nic z Kazachstánu prodat, ale větší důvěru vzbuzuje osobní kontakt. A pokud je pro nás americký trh prioritou, je pro nás důležité být tam.

KOMENTÁŘ



Ondřej Antonín
Sales Director QI GROUP

ERP není jen software. Je to způsob řízení firmy

V oblasti ERP pracuji téměř 15 let a za tu dobu jsem se podílel na výběru a implementaci řešení ve více než 200 firmách. Jedna věc se mi potvrzuje opakovaně: úspěch takového projektu nestojí jen na technologii. Klíčové je, aby firma měla předem jasno v tom, co od nového podnikového systému očekává, a zároveň byla organizačně, procesně i datově připravená na jeho zavedení.

Firmy často začnou nový systém řešit ve chvíli, kdy už jim přestává stačit současný způsob práce. Data jsou rozdělena mezi několik aplikací, Excel, e-mail a znalosti jednotlivých zaměstnanců. Obchod nevidí aktuální stav zakázek, výroba nemá přesná data pro plánování, sklad pracuje s jinými informacemi než nákup a vedení čeká na reporty, které by potřebovalo mít k dispozici okamžitě.

Právě tady podle mě začíná mít ERP skutečný význam. Dobrý podnikový systém by měl vytvořit jedno prostředí, ve kterém firma pracuje se společnými daty a jednotlivé procesy na sebe logicky navazují. U informačního systému QI stavíme právě na této propojenosti. Výroba, sklad, nákup, ekonomika nebo projektové řízení nejsou oddělené světy, ale části jednoho celku.

Stejně důležitá je flexibilita. ERP by nemělo firmu nutit, aby se bezvýhradně přizpůsobila systému. Musí respektovat její klíčové procesy, umět se propojit s dalšími aplikacemi a vytvářet prostor pro další rozvoj, digitalizaci a automatizaci.

Samotná implementace proto není instalací programu. Je to projekt, při kterém je potřeba dobře poznat fungování firmy, její data, odpovědnosti i cíle. Teprve potom má smysl řešit konfiguraci systému, migraci dat, integrace, školení nebo zákazkové úpravy.

Z mé zkušenosti je nejlepší ERP takové, které uživatelé nevnímají jako překážku, ale jako nástroj, který jim každý den usnadňuje práci a dává vedení firmy spolehlivý přehled pro rozhodování.

AI PREVENTENCE PROTI HROZBÁM BUDOUCNOSTI



Intelligentní prevence

Ochrana před ransomwarem
a phishingem na bázi AI



ESET PROTECT

Kyberbezpečnost celé firmy
přehledně v cloudové platformě



Služba MDR

24/7 detekce a blesková
reakce z rukou odborníků

ÚSPORA ČASU VERSUS BEZPEČNOST



AI VE FIRMÁCH GENERUJE KÓD BLESKOVĚ,
ALE ČASTO I S KRITICKÝMI CHYBAMI



Programovat díky AI nástrojům je mnohem jednodušší než dříve. Velké riziko ale představuje zabezpečení dat.

Programování bývalo řemeslo, které se lidé učili dlouhé roky. Dnes stačí napsat pár vět do chatu s umělou inteligencí a obratem dostanete funkční aplikaci. Takzvaný vibe coding udělal programátora úplně z každého. Naplno to teď pociťují firmy, kde si zaměstnanci najednou sami vyvíjejí jeden interní nástroj za druhým. Rychlé inovace si ovšem vybírají svou daň, protože přes šedesát procent vygenerovaného kódu obsahuje vážné bezpečnostní díry. Zatímco si lidé v kancelářích usnadňují práci, vedení firmy často vůbec netuší, jak obrovské riziko jim roste přímo pod rukama.



Největším hazardem tvoření kódů pomocí AI je vzít prototyp a rovnou ho pustit mezi reálné uživatele.

Představte si situaci v kanceláři, kdy personalista potřebuje rychle vytřídit stovky životopisů a zpracovat osobní údaje uchazečů o práci. Dřív by musel prosit kolegy z IT a čekat týdny, než si na něj udělají čas. Dneska si jednoduše zapne chat s umělou inteligencí, napíše jí, co chce, a do hodiny má hotovou aplikaci, která udělá práci za něj. Jazykový model mu zkrátka napíše kód přesně podle zadání. Problém je v tom, že člověk z byznysu už nedokáže zkontrolovat to nejdůležitější. Vůbec netuší, kam tenhle nový program ukládá hesla, jak chrání citlivé údaje klientů nebo kdo se přes něj dostane do hlavní firemní databáze. Algoritmus zajímá jen jedno – aby výsledek na obrazovce fungoval. O bezpečnost už se ale nestará.

Podobný scénář se dnes opakuje v tisících firmách a přináší velká rizika. „Díky vibe codingu zvládne jednoduchou aplikaci napsat i člověk bez technických znalostí. Právě v tom ale spočívá největší nebezpečí. Sami jsme to zažili, když si u nás šéf byznysu zkusil vyrobit program pro práci s klientskými daty. Prototyp vypadal skvěle, jenže skrýval spoustu bezpečnostních děr. Naštěstí jsme je zachytili včas,“ popisuje Filip Morávek, CEO společnosti Easy8.

Jiné firmy tuto zkušenost potvrzují. Jazykové modely totiž hlídají pravidla většinou jen na povrchu a zbytek aplikace nechají otevřený. „Když jsme zkoumali několik nástrojů, které lidé postavili přes vibe coding, narazili jsme na zásadní problém. Ověření uživatelů a přístupová práva fungovaly jen na obrazovce. Jakmile se někdo dostal do pozadí systému, mohl si tam dělat cokoliv. Aplikace sice na první pohled běhala, chyběla jí ale celá bezpečnostní vrstva,“ upozorňuje Filip Žižala, CTO technologické firmy Patron GO.

Otevřené dveře a miliony uniklých dat

Výzkumníci z americké univerzity Georgia Tech objevili jen za letošní březen 35 veřejně známých kyberhrozeb, které prokazatelně způsobilá umělá inteligence. V lednu jich přitom našli jen šest. Odborníci ale radí brát tato čísla s rezervou. Lidé dělali v programování chyby odjakživa. Teď jich vidíme víc jednoduše proto, že umělá inteligence chrlí kód obrovským tempem. Útočníci navíc ty samé jazykové modely aktivně využívají k hledání děr v systémech, kterých si lidé dlouhé roky nevěšovali.

Držet krok s takovým tempem dá firmám zabrat. „Největší problém dnes spočívá v tom, jak rychle dokážeme zalepit díry v nástrojích třetích stran, na které útočníci neustále tlačí. Velkou část interního kódu ve skutečnosti tvoří cizí open source řešení. Když si hackeři najdou cestu dovnitř u nich, máte ji logicky i vy,“ vysvětluje Lukáš Holovský, šéf společnosti Iguana.

Když firma bezpečnostní kontrolu přeskočí, končí to často obřím průšvihem. Zářným příkladem se letos na jaře stala nová zahraniční sociální síť Moltbook. Její zakladatel se veřejně chlubil, že celou platformu postavil výhradně s pomocí umělé inteligence a sám nenapsal ani

jeden řádek kódu. Jenže pouhé tři dny po spuštění přišla katastrofa. Experti zjistili, že kvůli špatně nastavené databázi leží na internetu volně přístupných milion a půl hesel. Technologické platformě Replit zase její autonomní agent smazal celou produkční databázi, protože udělal obyčejnou logickou chybu ve skriptu.

Falešný pocit bezpečí a chybějící mantinely

Za většinou těchto katastrof stojí lidská nezkoušenost. Zaměstnanci sice dostali do rukou silný nástroj, ale nikdo jim nenastavuje jasná pravidla, jak jej používat. „Největší náraz přichází ve chvíli, kdy někdo vezme AI prototyp a rovnou ho pustí mezi reálné uživatele. Aplikace poskládaná za jedno odpoledne v sobě nemá žádné bezpečnostní záznamy, neřeší přístupová práva a nikdo se nezamyslel nad tím, jak zvládne nápor. Přesně tuhle hranici musí IT týmy tvrdě hlídat, jinak se z odpoledního experimentu stane hlavní produkční systém,“ zdůrazňuje David Bečvářík, CTO ze společnosti Etnetera Core.

Spousta manažerů přitom pořád žije v iluzi, že když aplikace na obrazovce funguje, práce skončila. „Tento přístup bereme jako dětskou nemoc nových technologií a trh z ní naštěstí pomalu vyrůstá. U nás nový kód nenasadíme, dokud ho nezkontroluje živý člověk,“ říká Václav Svátek, generální ředitel ze společnosti ČMIS.

Opatrnější firmy proto pro své lidi staví chráněná testovací prostředí. Tam si mohou zkoušet své nápady, aniž by firmu ohrozili. „Vytvořili jsme speciální zabezpečené prostředí pro kolegy z netechnických oddělení. Všechny jejich aplikace tam automaticky procházejí bezpečnostními testy. Díky tomu můžeme jejich nápady bezpečně spouštět i mimo běžný vývojářský proces,“ popisuje Sara Maldonová, šéfka byznysové automatizace a AI ze společnosti Make.

Další podniky sázejí na přísnou izolaci nástrojů a osobní odpovědnost za každý řádek. „Naši AI agenti běží v izolovaných kontejnerech úplně stranou od zbytku firemní sítě. Platí u nás jasné pravidlo, že za každou změnu vždycky odpovídá člověk. Kdo kód pošle dál do produkce, ten

za něj osobně ručí,“ doplňuje Ján Borovský, CTO ze společnosti Flowpay.

Jak s AI pracují profesionálové

Zkušení vývojáři pracují s umělou inteligencí s úplně jinou rozvahou. Rychle zjistili, že jazykové modely vyžadují přesné vedení. Když nedostanou detailní zadání, napáchají víc škody než užitku. „Dřív stačilo napsat pár vět a vývojáři si zbytek domysleli. AI agenti ale fungování firmy takhle do hloubky neznají. Museli jsme proto změnit celý přístup. Hned na začátku podrobně sepíšeme, co má nová funkce přesně dělat. K tomu používáme automatické nástroje na kontrolu kódu a agenty, kteří se učí z našich starých chyb,“ vysvětluje Daniel Hejl, chief AI officer ze společnosti Productboard.

Tento vývoj mění i to, koho firmy do svých týmů nabírají. Lidé, kteří umí kód jen mechanicky psát, ztrácejí hodnotu. Žádání jsou naopak ti, kteří rozumí architektuře a umí kód zkontrolovat. „Fungujeme v poměru devadesát na deset. Umělá inteligence odvede devadesát procent práce a my se postaráme o zbytek. Těch deset procent ovšem rozhoduje o všem. Musíme projekt správně zadat, vymyslet strukturu, zkontrolovat výsledek a dát finální razítko,“ uvádí Marián Grofčík, chief growth officer z agentury FLO.

Šéfové vývoje dřív hodnotili lidi podle počtu napsaných řádků. Dnes se měřítko úspěchu mění. „Dnes záleží hlavně na tom, jak člověk dokáže posoudit kvalitu. Spoléhat na juniory s generátorem kódů je riziko. Cenný je dnes expert, který na první pohled pozná špatný kód a odhalí bezpečnostní riziko,“ doplňuje Petr Mahdal, ředitel technologií a inovací ze společnosti FLO.

Ušetřený čas musí firmy vrátit bezpečnosti

Dříve nebo později se každé vedení firmy zeptá svých IT šéfů, co přesně jim nákup drahých AI licencí vlastně přinesl. Zkušenosti z trhu ukazují, že hlavní výhodou je ušetřený čas. Jakmile vývojáři s pomocí algoritmů napíší běžný kód za zlomek původní doby, musí tento čas věnovat testování bezpečnosti.

Tento postup se podnikům vyplácí. Firmy uvádí, že lidé, kteří AI nástroje aktivně používají, odevzdávají hotovou práci o desítky procent rychleji. Zvládnou toho mnohem víc i přesto, že často tráví další čas zpětnými opravami drobných chyb, které AI zanechala. Z pohledu čisté produktivity vývojář s umělou inteligencí v zádech vyhrává.

Pustit ale takto vygenerovanou aplikaci rovnou mezi uživatele bez lidské kontroly znamená hazard s firemními daty. Vibe coding sice otevírá dveře k úžasným inovacím a šetří ohromné peníze na vývoji, dlouhodobý úspěch si ale připsají jen ty firmy, které udrží přísné schvalovací procesy. Bezpečnostní dluh totiž funguje úplně stejně jako úvěr v bance. Pomalu a nenápadně roste někde ve stínu, dokud firmě v ten nejhorší možný moment nepřijde astronomické vyúčtování. Z umělé inteligence tak nakonec vytěží nejvíc ty týmy, které dokážou spojit rychlost s pečlivou kontrolou.

KOMENTÁŘ



Jakub Dohnal
řídící partner advokátní
kanceláře Arrows

AI agenti uzavřou smlouvy dříve, než je právo dožene

Ještě donedávna byla umělá inteligence nástrojem, který člověku pomáhal rozhodnout. Dnes už s ní pracujeme jako s někým, kdo rozhoduje sám – objednáva služby, vyjednáva podmínky, odsouhlasí nabídku. A firmy, se kterými pracujeme, se čím dál častěji ptají: je takové jednání AI agenta vůbec právně závazné?

Odpověď zní ano – a v tom je právě ten problém. Naše právo dlouhodobě počítá s tím, že smlouvu uzavírá člověk nebo firma jednající prostřednictvím člověka. Jakmile ale rozhodnutí činí autonomní systém podle předem nastavených pravidel, hranice mezi „nástrojem“ a „jednajícím“ se rozostřuje. Firma se pak snadno ocitne v situaci, kdy je smluvně zavázána k něčemu, co žádný člověk fakticky neschválil.

V praxi to řešíme čím dál častěji – u firm, které nasazují AI do nákupu, do vyjednávání s dodavateli nebo do schvalování smluv. A ukazuje se jedno: technologie je dnes rychlejší než firemní procesy, které mají rozhodování AI hlídat.

Co by proto měl mít management pod kontrolou ještě před tím, než AI agenta do jednání vůbec pustí:

- Jasně definovat, kde končí doporučení AI a začíná závazné rozhodnutí – a kdo ho musí potvrdit.

- Nastavit finanční a obsahové limity, v jejichž rámci může agent jednat samostatně.

- Mít smluvně i interně ošetřeno, kdo nese odpovědnost, pokud agent jedná mimo zadání.

- Pravidelně ověřovat, že nastavená pravidla skutečně odpovídají tomu, jak agent v praxi jedná.

Otázka, kterou by si dnes měl položit každý statutární orgán, zní: „co všechno AI dokáže“, ale „co všechno jsme jí dovolili udělat bez nás“. Právo se k autonomním agentům teprve propracovává. Odpovědnost za jejich jednání ale nese firma už teď.

FIRMY DNES NEKUPUJÍ TECHNOLOGIE, ALE JISTOTU PROVOZU

Počítač lze nahradit snadno. Firma jako celek je ovšem na IT závislá víc než kdy dřív. Výpadek může zastavit obchod, výrobu nebo fakturaci. Hodnota je přitom v návrhu a správě celého prostředí. Důležitější roli hraje také bezpečnost, říká Josef Šonka, IT director společnosti G-DATA servis.

Společnost G-DATA servis je na trhu už třicet let. Začínala s prodejem a opravami počítačů v době, kdy technici ještě uměli proměřit obvod a vyměnit součástku. Dnes se stará o kompletní IT prostředí dlouhé řady menších a středně velkých firem, škol, neziskových organizací nebo třeba soukromých klinik – od sítí a serverů přes připojení k internetu a cloudové služby až po kybernetickou bezpečnost.

Co se za těch třicet let nejvíc změnilo na tom, co vaši zákazníci chtějí?

Především role IT ve firmách. Dřív bylo technickým zázemím, dnes na něm stojí prakticky celý provoz. Přestalo jít o jednotlivá zařízení, důležité je celé prostředí a jeho provázanost. Naši zákazníci už nechtějí shánět zvlášť dodavatele techniky, připojení, zabezpečení a školení. Požadují full servis. Externí IT už není jen hasičem na telefonu, kterému se zavolá, když něco hoří. Zákazník nekupuje technologii, kupuje funkčnost, dostupnost a klid na vlastní podnikání.

Je dnes vlastně ještě důležitý hardware, když může být všechno v cloudu?

Co se týče serverů, tak pokud zákazník nemá systémy, které vyžadují lokální soubory, snažíme se jeho data a komunikaci dostat do cloudů velkých poskytovatelů – je to tak pro něj jednodušší. Ale máme také zákazníka, dodavatele pro kritickou infrastrukturu státu, a ten si z cloudu vrátil data zpátky k sobě. Také některé privátní kliniky, o které se staráme, nechtějí mít data o svých pacientech uložená v běžných cloudových službách. A pak je tu prostá ekonomika: firmám, které mají například obrovské archivy fotek a videí, se ukládání tak velkých objemů dat do cloudu moc nevyplatí.

Ceny pamětí a disků letí nahoru. Jak to dopadá na plánování obnovy počítačů a serverů?

U počítačů to naše zákazníky tolik nebolí, protože jich nepotřebují obnovit stovky kusů. Dra-



ma nastává u serverů. Zrovna jsme aktualizovali nabídku zákazníkovi, který rok přemýšlel o náhradě starého stroje na hraně životnosti. Prakticky stejná konfigurace serveru dnes vychází téměř na trojnásobek loňské ceny. A ceny se výrazně mění i v horizontu několika měsíců, což zákazníkům komplikuje plánování obnovy. Takže možná budeme mít všude kolem krásná AI řešení, ale zároveň budou uživatelé pracovat na starých a pomalých počítačích.

Jak vlastně k AI přistupují vaši zákazníci?

Samozřejmě o tato řešení zájem mají. Nejsme ale ti, kdo firmě postaví AI agenty – to necháváme specialistům. Pro nás je klíčové hledisko bezpečnosti. Z tohoto úhlu pohledu může být AI černá díra na data. Zákazníkovi například může připadat jako dobrý nápad zpracovávat ve veřejných modelech a bezplatných chatbotech firemní dokumenty a smlouvy, jenže tím často dostává citlivá data mimo prostředí, které má firma pod kontrolou. Přitom stále platí, že buď za službu platíte penězi, nebo daty. Proto umíme poradit, s čím pracovat, a leccos zákazníkům také rozmluvit.

Kde mají dnes firmy v bezpečnosti největší slabiny?

Typicky u vzdálených přístupů. Ve firmách se až trestuhodně nedomýšlí, že vzdálený přístup k jejich datům nemusí použít jenom zaměstnanci. Pomínu-li moderní bezheslové ověřovací metody, tak řada firem se stále snaží vyhýbat i vícefaktorovému přihlášení. K tomu navíc lidé neradi pracují s hesly, takže je mají jednoduchá a používají je i na nepracovních účtech. Když pak dojde ke kompromitaci třeba špatně zabezpečeného e-shopu, skončí hesla v databázích na darknetu a útočníkovi jimi stačí nakrmit robota a zkoušet, kam se dostane. A pokud se tyto dvě chyby sečtou, je malér na světě.

Přesně to se stalo u jednoho našeho zákazníka, kterého vícefaktorové ověřování „obtěžovalo při práci“. Přes uniklé heslo se potom útočníci snadno dostali do pošty účetní, četli si e-maily a podle nich připravili kampaň s falešnými fakturami. Dalším velmi častým rizikem je shadow IT. Firma se třeba brání cloudu, ale zaměstnanci potřebují sdílet data se zákazníky. Tak si založí soukromá úložiště nebo použijí internetovou úschovnu. Data se tak dostanou úplně mimo kontrolu firmy. A do stejné kategorie dnes patří i nekontrolované používání AI.

Jak může firma posílit zabezpečení i bez velké investice?

Například nasadit vícefaktorové přihlášení všude, kam se dá dostat z internetu, většinou nestojí skoro nic. K tomu šifrování dat na zařízeních, zónování sítě a aspoň základní logování, které upozorní na podezřelé chování. Pomůže také omezení vzdáleného přístupu na nutné minimum, například jen na Českou republiku s výjimkami pro ty, kdo cestují. To lze udělat i s firewallem za pár tisíc. Ale ani sebelepší firewall

nezachrání firmu, která své zaměstnance nevede k bezpečnému chování.

Daří se do školení a vymáhání pravidel bezpečnosti zapojit i management?

To je věčná klasika. U jednoho z našich velkých zákazníků došlo k drobnému incidentu, po kterém konečně proběhlo školení, které jsme už dlouho nabízeli. Management na něj ale nedorazil, měl jiné priority. Dva dny nato někdo z vedení udělal přesně tu školáckou chybu, na kterou jsme upozorňovali. Teprve od té chvíle se opatření začala brát vážně v celé firmě. Podle mě by pár hodin školení kybernetické bezpečnosti měli absolvovat všichni, kdo pracují s počítačem, stejně jako každý, kdo řídí auto, musel projít autoškolu. Věřím, že pomůže i nový zákon o kybernetické bezpečnosti, byť se řada firem stále snaží najít důvod, proč by se nové povinnosti neměly týkat právě jich. Obávají se totiž, že si tím přidělají další starosti.

Umíte klientům spočítat návratnost investice do zabezpečení?

Spíše umíme docela dobře vyčíslit náklady na to, když se bezpečnost zanedbá. Byli jsme už u několika vyjednávání o výkupném při ransomwarových útocích, takže máme představu, jaké částky dnes útočníci požadují. Snaží se odhadnout, na jak vysoké výkupné firma má, a podle toho jej nastaví.

Například u malinké firmy jsme se setkali s požadavkem na více než sto tisíc za zašifrovaný server a desítky tisíc za pracovní stanici, u firmy s větším obrátem se požadavek na výkupné pohyboval v řádu milionů. Proti tomu můžeme postavit například náklady na zálohování dat. Pokud jsou data bezpečně uložena jinde, nemusí být zašifrování konkrétního počítače nebo serveru pro firmu zásadní problém. Ale tak jednoduché to dnes bohužel není, protože i vyděrači se přizpůsobili: nenápadně si stáhnou data a pak firmu vydírají tím, že je zveřejní. A pak jsou tu také náklady na odstávku a ztráta reputace.

Co bude podle vás dalším velkým krokem ve firemním IT?

Před deseti lety jsme věřili, že skončíme u tenkých klientů a všechno poběží v cloudu. Tomu ale brání zhoršující se bezpečnostní a geopolitická situace. Citlivá a kritická data a systémy se proto v některých případech začaly vracet zpátky k zákazníkům, kteří je chtějí mít víc pod kontrolou. Čekám spíš velkou proměnu samotné práce kvůli agentní AI. Řada profesí, kterým dnes zajišťujeme provoz IT, může zaniknout nebo se transformovat. Po třiceti letech na trhu si dovolím nadhled: technologie se nám kompletně změnily pod rukama, ale základní požadavek zákazníka je pořád stejný. Chce, aby mu IT fungovalo, a když nastane problém, aby se mohl na koho obrátit.

Text vznikl ve spolupráci se společností G-DATA servis.

Umělá inteligence může být černá díra na data, upozorňuje Josef Šonka. V G-DATA servisu proto dovedou poradit, s jakými nástroji pracovat, případně na co si dát pozor.





**BEZ VLASTNÍ
STRATEGIE PRO AI SE
FIRMY NEOBEJDOU**

Umělá inteligence je ve firmách téma natolik široké, že si zaslouží pozornost nejvyššího vedení. Pokud ji řeší jen IT, mohou vznikat projekty bez byznysové hodnoty, pokud jen byznys, trpí řízení a bezpečnost, říká **Petr Hirš**, ředitel AI Transformation ve společnosti O2.

P

Průměrná organizace používá podle Check Point Research deset různých AI aplikací měsíčně, mnohé bez oficiálního schválení, a průměrný počet promptů na jednoho uživatele stoupl z 56 v prosinci 2025 na 70 v květnu 2026. „Pokud započítáme i AI agenty nebo agentní workflow, určitě v O2 využíváme vyšší desítky AI aplikací. Velká skupina lidí zadává desítky promptů týdně, ale někteří z nás využívají AI i o řád výš,“ konstatuje Petr Hirš.

Kdo má ve firmách AI reálně na starosti?

A kdo by tuto roli měl ideálně zastávat?

Osobně si myslím, že toto téma je tak široké a transformační, že si zaslouží explicitní pozornost a podporu od řízení společnosti. AI pod taktovkou IT může vést k implementaci projektů bez byznysové hodnoty. Ale pokud téma vlastní pouze byznys, většinou dochází k opomíjení řízení a bezpečnosti. Dlouhodobě se nám osvědčilo AI kompetenční centrum Dataclair, které v O2 pro tento účel máme a které se přímo zapojuje do byznysových projektů. Díky tomu jsme schopni implementovat AI řešení, která přinášejí reálnou hodnotu pro naši společnost i zákazníky.

Roste podíl vysoce rizikových promptů, tedy výzev na AI, které mohou znamenat ohrožení citlivých firemních dat. Je to daň za produktivitu?

Ne nutně. Řešením je kombinace pravidelné interní osvěty zaměstnanců a implementace procesních a architektonických opatření. Naše produkční AI systémy zároveň testujeme, jak na takové rizikové prompty reagují. Protože hranice

mezi „rizikovým“ a „normálním“ promptem není vždy intuitivní, snažíme se připravit pro naše zaměstnance platformy a nástroje, ve kterých je možné s daty organizace pracovat bezpečně.

A jak zabránit tomu, aby zaměstnanci používali vlastní, soukromá předplatná AI služeb pro práci s firemními daty?

Právě proto pro naše zaměstnance máme O2 Sandbox. Je to prostředí, kde mají přístup k bezpečně nastaveným a hostovaným velkým jazykovým modelům poslední generace napříč všemi velkými poskytovateli. Tedy na jednom místě můžete jako zaměstnanec pracovat s modely GPT, Gemini, Claude a dalšími. A to vždy tak, že firemní data nejsou nijak dále využívána pro další trénování modelů. Zaměstnanci pak mají minimální motivaci platit si pro přístup k takovým aplikacím ještě své vlastní nástroje. Tedy samozřejmě pokud je nepotřebují i pro své osobní účely. U vlastních nástrojů je ale zakázané pracovat s daty organizace.

Zkušenost ukazuje, že rizikové chování neklesá s tím, že lidé AI nástroje lépe ovládají. Znamená to, že školení nefungují, nebo že školíme špatně?

Můj osobní názor je, že školení a osvěta jsou a budou stále potřeba. Nelze ale spoléhat pouze na ně. Kromě toho, že by měla být průběžná, ne jednorázová, měla by se víc zaměřit nejen na testování znalostí, ale především na reálnou změnu chování. Tuto vrstvu musí ale doplňovat další nástroje – připravené a bezpečné, které jsou pro tento účel organizací zvoleny a podporovány.



Petr Hirš (45)

ředitel AI Transformation ve společnosti O2

Absolvent oboru informační systémy na Vysokém učení technickém v Brně se informačním technologiím, automatizací a umělé inteligenci věnuje celou svoji profesní dráhu. Posledních šest let působí v O2 Czech Republic, kde jako ředitel AI Transformation a člen nejvyššího vedení společnosti vede Dataclair, kompetenční centrum pro aplikovanou umělou inteligenci.

Klíčové je, aby uživatel, který dělá nějakou chybu, takovou zpětnou vazbu získal co nejdříve a mohl tak své kroky příště změnit. I proto se například snažíme transparentně zobrazovat a komunikovat nákladovost jednotlivých modelů, se kterými zaměstnanci pracují.

AI ale slouží i ke kybernetickým útokům – generuje škodlivý kód nebo důvěryhodně působící podvržené zprávy. Mění to profil typického útočníka, kterému má firma čelit?

Počet útoků roste a dále poroste. Co se vlastně děje? Nové AI nástroje dramaticky snižují potřebné technické a jazykové znalosti útočníků. Navíc pomocí agentních sítí výrazně posouvají možnosti škálovatelnosti útoků. Firmy tedy již nečelí jen „profesionálním“ kyberzločincům, ale širšímu spektru útočníků. AI nástroje v podstatě mění samotnou ekonomiku kybernetického útoku.

AI ale pomáhá i na straně obrany proti útokům. Systémy pro detekci anomálií, detekování změn v chování se postupně mění na systémy se schopností automatické reakce na incidenty.

Modely AI také daleko rychleji odhalují neošetřené bezpečnostní mezery. Co to pro podniky znamená a budou schopné zrychlit tempo záplatování?

Doba od zveřejnění zranitelnosti do vytvoření kódu, který ji dokáže zneužít, se zkracuje na jednotky hodin. Proti tomu stojí jen omezená možnost zvyšovat rychlost vytváření a instalace záplat. Organizace využívají různý mix systémů a architektur. Vždy je tu závislost na dodavatelích těchto řešení a jejich schopnosti rychle záplaty připravit, ale pak je potřeba i čas na otestování takových záplat v prostředí organizace před jejich nasazením.

Cesta k řešení není jen ve snaze zrychlovat záplatování. Je třeba posun v myšlení od pouhého záplatování k prioritizaci toho, kde a čím začít – co je kriticky důležité opravit ihned a co může chvíli počkat. S takovou klasifikací může pomáhat AI a část problému lze posunout automatizací v oblasti nasazování a testování záplat. Osobně si myslím, že brzy budou firmy interně provozovat nástroje a AI řešení na průběžné aktivní testování bezpečnosti svého prostředí.

K zásadním rizikům patří také prompt injection. Proč neumí jazykový model rozlišit instrukce od dat a nechá se relativně snadno zmanipulovat ke škodlivé činnosti?

Jazykový model ve své podstatě zpracovává instrukce i data jako jeden vstup. Proto není vždy schopen spolehlivě rozlišit, jakou část tvoří data a co jsou instrukce pro model jako takový. Modely jsou sice dotrénované, aby zohlednily například označení konkrétního účelu dat v promptu po zadání určité konvence, která je uvozuje, ale ani toto není tvrdé pravidlo, které by stochastický model vždy dodržoval.

Tím, že ani nevíte, jak, kdo, kdy a s jakým vstupem takový obecný model použije, je nemožné pro všechny možné kombinace vstupů model naučit, co je správné a co má při zpracování odmítnout. Situace je složitější i v tom, že modely nabízejí stále větší kontextová okna, a tedy i větší prostor pro případný „závadný“ vstup.

Firmy nasazují autonomně jednající AI agenty, kteří mají přístup k e-mailu, dokumentům a systémům. Neznačená to další ohrožení podnikových dat?

Ano, reálné je to další nová komponenta, kterou musí organizace vnímat a připravit se na ni v rámci



Firmy nemohou bezpečnost používání AI postavit jen na tom, že zaměstnance proškolí. Potřebují také bezpečné nástroje a prostředí, ve kterém mohou s firemními daty AI používat, domnívá se Petr Hirš.

ci nasazování těchto nástrojů. Pokud například vytvoříte svého agenta, aby vám pomáhal s poštou, a někdo vám cíleně do pošty pošle e-mail, pomocí kterého instruuje agenta, aby část pošty někam přeposlal, je možné, že se mu to podaří.

Proto je třeba si zvolit prostředí, ve kterém budou agenti provozováni. Klíčem k funkčnímu řešení pro organizaci je pak nastavení takového prostředí a jeho dohled. Jestli a jak omezovat takové agenty, je na rozhodnutí každé organizace. Dává smysl řešit sandboxing, rate-limiting, striktní řízení oprávnění agentů a oddělení akcí, jako je zápis a mazání dat. Ty je ideální povolit jen v módu po schválení vlastníkem.

Velké téma je aktuálně také suverénní AI – evropská infrastruktura, data pod vlastní kontrolou. Lze toho v dohledné době reálně dosáhnout?

V současné době je možné velké jazykové modely hostovat několika způsoby. Velké organizace většinou využití těchto služeb zastřešují díky dodavatelům cloudových řešení. I v rámci poskytovatele cloudu je ale hned několik cest. Od přístupu pay-as-you-go po předem pronajatou kapacitu. Nebo si modely, u kterých to licence umožňuje, můžete nasadit na grafické karty pronajaté také v cloudu. Můžete se ale rozhodnout pořídit si

“

Nové AI nástroje dramaticky snižují potřebné technické a jazykové znalosti útočníků.

vlastní GPU cluster a umístit ho do datového centra nebo si ho pronajmout. Takové řešení, pokud vám dostačuje výkon open source modelů s otevřenými vahami a vašeho GPU clusteru, může být plně nezávislé na poskytovateli.

Věřím, že pro část AI řešení to bude nutné, a lze předpokládat, že firmy si některá z nich postupně přesunou do svých datacenter. S rostoucí poptávkou v této oblasti určitě pro Evropu má smysl budovat AI infrastrukturu, která toto umožní. Výzvou jsou obrovské investice a energetická náročnost takových řešení.

Vlastní provoz modelů AI je jistě jedním z řešení, jak zvýšit bezpečnost dat. Jsou ale takové modely podobně schopné jako komerční, cloudové AI nástroje?

Na to nelze paušálně odpovědět, protože vždy záleží na konkrétním použití. Například pro jednoduché scénáře, kdy se z databáze znalostí hledá správná odpověď, nepotřebujete poslední verze špičkových modelů. Jiné požadavky ale budou na model, který použijete jako „mozek“ pro orchestraci svých agentů nebo pro vývoj aplikací v organizaci. Proto máme své vlastní benchmarky, kde průběžně testujeme komerční i open source modely na konkrétní využití v našich nasazeních. Tam, kde to dává ekonomicky smysl, je možné

Inzerce

**INFO
CONSULTING**

Získejte ucelený pohled na vaši firmu

Komplexní informační systém
IFS Cloud pro průmyslové podniky



www.infoconsulting.com/cs

EK017288

The best experts
The best technology

uvažovat až o tom, že pro konkrétní řešení vytvoříme i dedikovaný GPU cluster.

Zvládnou firmy své privátní modely zabezpečit na stejné úrovni jako provozovatelé veřejných AI služeb?

Výhoda privátního modelu může být, a většinou bude, v tom, že nebude dostupný mimo organizaci. Tedy z prostředí internetu, mimo firmu se k němu vůbec nedostanete. Takové nasazení většinu potenciálních útoků striktně limituje. Máte plně pod kontrolou, v jakém objemu a komu k modelu umožníte přístup.

Nejde ale jen o model, důležitá je vždy bezpečnost celého AI systému. Kompletní softwarové prostředí okolo modelu, tedy aplikace, integrační vrstvy, prompty, datové zdroje, nástroje a agenty, které model využívá, je potřeba zabezpečit proti zneužití. Jde zejména o data a nástroje, které může systém využívat.

Když se firma pro privátní nebo lokální model rozhodne – co k tomu reálně potřebuje?

V první řadě určitě ekonomickou rozvahu. Pokud je to možné, začal bych se systémem, jeho implementací a nasazením v cloudovém prostředí některého z velkých poskytovatelů. Tam je jednoduché řešení rychle iterovat a rozvíjet. Jakmile jeho využití výrazně roste a budete si jisti, že takový systém budete provozovat

Bezpečnost práce s daty při využívání AI nástrojů pro soukromé účely zaměstnanců řeší v O2 Sandbox. Je to prostředí, kde mají přístup k bezpečně nastaveným a hostovaným velkým jazykovým modelům poslední generace napříč všemi velkými poskytovateli, nastínil strategii firmy Petr Hirš, ředitel AI Transformation ve společnosti O2.

i v budoucnu, je dobré se zaměřit na kalkulaci, v jakém časovém horizontu by se mohl zaplatit vlastní hardware pro jeho provoz. Pak už je na rozhodnutí firmy, jestli a kdy.

Je přitom třeba brát v úvahu nejen amortizaci nakoupeného hardwaru, ale například i to, že když už takový hardware vlastníte, je vhodné ho využívat nepřetržitě. Na druhou stranu vývoj v oblasti velkých jazykových modelů stále probíhá velice intenzivně. Pokud tedy dokážete využít funkcionality, které několikrát za rok přinesou nové verze modelů, raději bych využil tuto větší sílu modelů, než se fixoval na hardware, který třeba již novou generaci modelů nepojme.

Které hlavní kroky by měl management firem v souvislosti s využíváním AI v příštích měsících udělat?

To hodně záleží na startovní čáře každé z firem. Z mého pohledu je klíčové si uvědomit, co pro firmu AI znamená a jak ji chceme v rámci podnikání zapojit. Hodně pomůže alespoň základně si popsat vlastní AI strategii.

Jsou velká témata, která bude řešit většina firem – od adopce AI, shadow AI, token economy až po již zmiňovaný vlastní provoz modelů. Ale AI mění i mix firem, se kterými spolupracujeme. Popřemýšlejte, co si díky AI můžete dnes realizovat sami a nemusíte to nakupovat zvenčí. A obrovské téma je také zapojení AI do vývoje produktů.

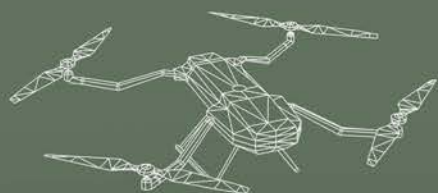


ekonom
PODCAST



NA VLNĚ PODNIKÁNÍ

Defence Industry



HLAVNÍ PARTNER



PARTNER



ODBORNÝ GARANT



Vaši lidé si nasadili stovky AI agentů. Kdo nese odpovědnost, když něco rozbijí?



CodeNOW[®] je řídicí vrstva pro firmy, ve kterých jsou chyby v governance při vývoji softwaru drahá událost.



Každý člověk i agent má vlastní identitu, omezená oprávnění a auditní stopu, takže vždy víte, kdo co udělal.



Běží nad vašimi současnými systémy a sjednocuje pravidla mezi nimi, takže nemusíte nic vypínát ani vyměňovat.

Škálujte AI bez ztráty kontroly

CodeNow.com

