

**PROCESSING INFORMATION
SECURITY REQUIREMENTS
FOR IT SYSTEMS
USING RAG-SUPPORTED
LLM****IT RENDSZEREKKEL KAPCSOLATOS
INFORMÁCIÓBIZTONSÁGI
KÖVETELMÉNYEK FELDOLGOZÁSA
RAG TÁMOGATOTT LLM
SEGÍTSÉGÉVEL**FEHÉR Dávid János¹**Abstract**

Managing information security requirements during IT system design poses significant challenges, as organizations must comply not only with international standards such as NIST 800-53 and ISO 27001:2022 but also with their internal regulations. Large Language Models (LLMs), particularly with support from Retrieval Augmented Generation (RAG), offer new possibilities for more efficient processing of these requirements. This publication highlights how LLMs provide real-time access to relevant information, reducing the time spent on manual research. While LLMs offer numerous benefits, their use also entails risks, such as potential misinterpretations. The research sheds light on key aspects of designing LLM-based systems and underscores the importance of human expertise in enhancing the efficiency and accuracy of managing information security compliance.

Keywords

Information security, LLM, RAG, Security compliance, Security control

Absztrakt

Az információbiztonsági követelmények kezelése az IT-rendszerek tervezése során komoly kihívást jelent, mivel a szervezeteknek nemcsak a nemzetközi szabványoknak, mint a NIST 800-53 és az ISO 27001:2022, hanem belső szabályzataiknak is meg kell felelniük. A nagy nyelvi modellek (Large Language Modell – LLM), különösen az adatlekérésre alapozott generálás (Retrieval Augmented Generation - RAG) támogatással, új lehetőségeket teremtenek a követelmények hatékonyabb feldolgozásában. A publikáció bemutatja, hogyan biztosítanak az LLM-ek valós idejű hozzáférést a releváns információkhoz, csökkentve a manuális kutatásra fordított időt. Az LLM-ek rengeteg pozitív lehetőséget rejtenek, de használatuk kockázatokat is rejt, például félreértelmezéseket. A kutatás rávilágít az LLM rendszerek tervezésének főbb pontjaira és az emberi szakértelem fontosságára, amely hatékonyabbá és pontosabbá teszi az információbiztonsági megfelelés kezelését.

Kulcsszavak

Információbiztonság, Nagy nyelvi modell, Adatlekérésre alapozott generálás, Biztonsági megfelelés, Biztonsági kontroll

¹ david.janos.feher@gmail.com | ORCID: 0000-0002-0742-8996 | PhD student, Óbuda University Doctoral School on Safety and Security Sciences | doktorandusz, Óbudai Egyetem Biztonságtudományi Doktori Iskola

BEVEZETÉS

IT vállalatok komplex környezetben komplex IT megoldásokat használnak. Ezen IT-megoldások tervezésében a rendszertervezők és mérnökök kulcsszerepet játszanak a megfelelőségi kontrollok technikai infrastruktúrákba történő beépítésében. Felelősségeik közé tartozik annak biztosítása, hogy az általuk választott, tervezett, kivitelezett, üzemeltetett és használt rendszerarchitektúrák, adatkezelési folyamatok, biztonsági protokollok és egyéb részletek összhangban legyenek a szervezet számára fontos szabványokkal, mint például a NIST 800-53 és az ISO 27001:2022 szabványokkal, de a kezelt adat és piaci terület függvényében területenként más és más követelménynek kell megfelelniük. A nemzetközi szabványokon kívül a szervezet saját működési módját, struktúráját és saját belső szabályzatait is figyelembe kell venni. Ezen követelmények nagy mértékben változhatnak a nemzetközi szabványok változása, szervezeti változások és rengeteg más egyéb esetben is, mind ezen információkkal naprakésznek lenni, minden dokumentumot elolvasni és észben tartani nagy terhet ró a szakemberekre. [1], [2], [3]

A mesterséges intelligencia azon belül is a nagy nyelvi modellek (Large Language Models - LLM) térnyerésével rengeteg területen próbálják ennek előnyeit használni, ide tartozik az információbiztonsági terület is. Az LLM jelentős potenciált rejt az információbiztonság területén lévő rengeteg követelményeket vagy ajánlásokat tartalmazó dokumentáció hatékonyabb feldolgozásához. Az LLM-ek még hatékonyabbak a Retrieval Augmented Generation (RAG) azaz visszakereséssel támogatott generálással kiegészítve, ez lehetővé teszi, hogy dokumentumokkal és egyéb fix információkkal még biztosabb válaszokat adjon a modell, sok esetben "grounding"-ként hivatkoznak erre. A mesterséges intelligencia területén a grounding azt jelenti, hogy egy modell vagy rendszer az általa használt információt külső valóságnak vélt adatokhoz kapcsolja, így biztosítva, hogy a generált válaszok pontosak és relevánsak legyenek. Ezen cikkben az információbiztonsági követelményekhez való hatékonyabb LLM-mel támogatott hozzáférés lehetőségét mutatom be. A követelmények könnyen hozzáférhetővé tétele nemcsak a hatékonyságot növeli, hanem ezen felül csökkentheti az esetleges félreértéseket és a mulasztások valószínűségét is a fejlesztési folyamat során.[4], [5], [6], [7]

LLM KÖVETELMÉNYEKET TARTALMAZÓ TUDÁSBÁZISHOZ

A kisebb és nagyobb IT-megoldások, platformok, folyamatok tervezésében és kivitelezésében kiemelten fontos a megfelelőségi szabványok szemelőt tartása, figyelembevétele, és lehetőség szerint a betartása. Ahhoz, hogy a rendszerek körülötte folyamatokba bevont IT-rendszertervezők, mérnökök, szakemberek meg tudjanak felelni ezen követelményeknek pontos és naprakész általuk feldolgozott és megértett információra van szükség. [1]

A technikai csapatok gyakran szembesülnek összetett megfelelőségi követelmény-nyel kapcsolatos kérdésekkel a rendszerek tervezési fázisa során. Mint például, hogy milyen követelményeknek kell megfelelnie a rendszerekben használt jelszavaknak. Mit és hogyan kell titkosítani, hogyan tudják biztosítani a megfelelő naplózását a rendszernek. Az LLM-ek részletes, kontextus függő válaszokat nyújthatnak ezekre a kérdésekre, kiterjedt tréningadataikból merítve az adott szervezeti környezetre szabott útmutatást. A nagy nyelvimodel-

lek ebben lehetnek segítségünkre, hiszen segítségükkel gyorsan lehet megválaszolni a technikai kérdéseket, ezzel valós idejű hozzáférést biztosítva a kontrollkövetelményekhez, annak leírásaihoz. [1], [8]

Az LLM-ek és az RAG szerepe a kontrollkövetelmények hatékony lekérdezésében

A ChatGPT és a Gemini, mint az AI ökoszisztéma szerves híres részei a “GenAI” azaz generatív mesterséges intelligencia jól ismert példái, amelyek LLM-ek azaz nagy nyelvi modellek nyelvfeldolgozási képességeiket kihasználva képesek „megérteni” a kontextust és emberihez hasonló szövegeket generálni, valamint összetett feladatokban, például szövegírásban és akár kontrollkövetelmények feldolgozásában is segíteni. A legnagyobb LLM-ek, nagyon kiterjedt korpuszon vannak tanítva, amely tartalmaz rengeteg az interneten elérhető forrást ebbe beleértve a legismertebb szabványokat is mint például a NIST 800-53 és az ISO 27001. Elmondható, hogy a ChatGPT és a Gemini mind a kettő kifejezetten jó válaszokat ad komplex kérdésekre, beleértve az információbiztonsággal kapcsolatos kérdésekre. Ezen tényezők lehetővé teszik, hogy azonnali, az LLM számára elérhetővé tett kontextus függvényében válaszokat adjanak információbiztonsági megfelelőséggel kapcsolatos kérdésekre. Például, ha egy szakember a NIST 800-53 egy adott kontrolljának implementálásához keres útmutatást, akkor az LLM alapú generatív mesterséges intelligencia rendszerek képesek használható magyarázatot, megvalósítási lépéseket és egyéb információkat nyújtani. Az egyik fő problémája az LLM modelleknek az, hogy félrevezető vagy hibás választ is adhat, de ezen probléma megoldására való a „grounding”, amely fogalomhoz tartozik az RAG architektúra. A rendszer rövid válaszára óriási segítség abban, hogy a tervezési döntésekhez vagy egyéb döntésekhez szükséges információk vagy ajánlások sokkal elérhetőbbek legyenek. A közel azonnali információhoz való hozzáférés révén csökken a komplex dokumentációk és leírások keresésével, olvasásával, értelmezésével töltött idő. Ez nem csak azt jelenti, hogy kevesebb időt töltenek ezzel az emberek, hanem hogy a kényelmes lekérdezésnek hála hamarabb fognak utána járni a dolgoknak. [1], [3], [6], [9]

TERVEZÉS ÉS MEGVALÓSÍTÁS EGYSZERŰSÍTÉSE LLM-EKKEL

A különböző informatikai rendszerek és platformok tervezésének előkészítése, tervezése, kivitelezése minden szervezet esetén és minden rendszer esetén valamennyire eltér, nagyban függ a szervezettől, a személyzettől és a tervezendő rendszertől. A fejezetben egy általános megoldást kivitelezését mutatom be, amelynek fő elemei könnyen átemelhetők az igények szerint. Az alábbi példák szemléltetik milyen típusú kérdésekre kereshetünk válaszokat a megoldással:

- Milyen titkosítás használható ügyfeladatok kezelésére?
- Milyen titkosítás használható személyes adatok kezelésére?
- Milyen követelményeknek kell megfelelni a jelszavak adásánál? Hány karakter hosszúnak kell minimum lennie, hogy megfeleljek a szervezet követelményeinek?
- Hogyan kell kezelni a szerverek a helyi adminisztratív felhasználójának jelszavát? Milyen sűrűn cserélendők ezeken a jelszavak?
- Milyen naplófájlokat kell gyűjteni és továbbítani? Hogyan?

Főbb lépések az LLM-ek folyamatokba történő integrálásához.

Minden szervezet esetén más az egyes lépéseknek a prioritása és a részletei és nagy mértékben függ az integrálás magától a folyamattól is, hiszen nem mindegy, hogy egy kisebb belső használatra szánt informatikai megoldásról beszélünk, vagy pedig egy igazán összetett szervezeti részlegeken átívelő komplex megoldásokat. Fontos azt is szem előtt tartani, hogy mely esetekben érdemes a folyamatot automatizálással támogatni, hiszen vannak annyira egyedi egyszeri esetek, amelyeket nem érdemes automatizálni. Ezen részletektől eltekintve általánosan az alábbi főbb lépéseket érdemes elvégezni az integráláshoz:[9], [10], [11]

1. **Megfelelő csapatok és szakemberek bevonása:** Az egyik legfontosabb lépés a szervezeten belüli megfelelő szakemberek megtalálása, hiszen a folyamat során rengeteg a szervezet meglévő információbiztonsági követelményeivel kapcsolatos kérdés merül fel. Érdemes bevonni mindenkit, kiemelten azokat, akik korábban kiemelt szerepet töltöttek be a hasonló kérdések megválaszolásában.[12], [13]
2. **Megfelelőségi követelmények és célok meghatározása:** Először egyértelműen körvonalazzuk az elemzés célját, minek akarunk megfelelni, mint például NIST 800-53 vagy az ISO 27001:2022 szabvány megfelelés, belső információbiztonsági követelményeknek való megfelelés, azon belül is tisztázni, hogy a konkrét részleg, vagy szervezet különböző követelményei miként milyen átfedéssel és prioritással van szemelőt tartva. Itt külön érdemes dokumentálni, amennyiben nem egységes a megfelelési követelmény a teljes szervezetben, hanem eltér megoldásonként vagy csapatonként. [12], [13]
3. **Megfelelőségi követelményekkel kapcsolatos dokumentációk gyűjtése:** Fontos minél korábban elkezdeni összegyűjteni, a megfelelőséggel kapcsolatban elérhető írott információkat és dokumentációkat. Amennyiben van dokumentáció, amely nem naprakész vagy hiányos vagy nem létezik, akkor annak a frissítését kezdeményezni kell vagy az eltéréseket dokumentálni.
4. **Megfelelő LLM-eszközök kiválasztása:** Olyan LLM mesterséges intelligencia modell és platformot válasszunk, amellyel meg tudunk felelni az általunk használni tervezett adatok és anyagok kezelésével kapcsolatos szervezeti adatbiztonsági követelményeinknek. A legnagyobb felhő szolgáltató cégek platformjain valószínűleg megtaláljuk a szükséges szolgáltatásokat, a kutatásban leírtak tesztelése és validálása során a Google Cloud Platform (GCP) Vertex AI nevű AI platform volt használva. AWS és Azure Cloud esetén is rendelkezésre állnak jelenleg a szükséges komponensek. [7], [14], [15]
5. **Megfelelőségi követelményekkel kapcsolatos írott anyagok biztosítása:** A rendszer csak a leírt adatokat tudja feldolgozni és figyelembe venni, abban az esetben, ha valamire nincs írott dokumentációnk, de szükségünk van az információ figyelembevételére, akkor gondoskodni kell ennek az információnak a leírásáról. Amennyiben van rá lehetőség akkor ezen dokumentumokat érdemes megszerezni és átnézni, hogy csak a valóban naprakész szabályzatok és részletek benne. Átfedés esetén priorizálni kell a dokumentumokat, például amennyiben van egy részlegszertinti követelmény, amely szigorúbb, mint a szervezet követelménye. A modern LLM modellek nagy token száma, ami azt jelenti, hogy nagy mennyiségű szöveg

- feldolgozására alkalmas, így ezeket kategóriákként össze lehet vetni, és ajánlásokat létrehozni a feltöltött dokumentumok kapcsán. [12], [13]
6. **RAG elő-feldolgozás:** Az összegyűjtött dokumentumok validálása után egy “embedding model” azaz beágyazási modell segítségével átalakítjuk. A beágyazási modellt esetünkben arra használjuk, hogy szövegeink feldarabolása után abból numerikus reprezentációkat hozzunk létre, így a szövegünk hatékonyan feldolgozhatók és összehasonlíthatók machine learning (gépi tanulási) modellekben. Ezen átalakítás után létrejött adatunkat egy vektor adatbázisban tároljuk. [7], [14], [15]
 7. **LLM alapú információfeldolgozás:** Lehetőségünk van az LLM megoldások segítségével is feldolgozni az adatainkat bizonyos szempontból, ezzel plusz mélységet adva a feldolgozott vektor adatainknak, akár metaadatként, akár egyéb kiegészítő információként.
 8. **Egyéb információ források integrálása:** Lehetőségünk van más egyéb információ forrásokat, egyéb forrásokat is tartalmazó adatbázisokat és API alapú információ forrásokat integrálni.[16]
 9. **Állandó promptok kidolgozása:** Az állandó promptok továbbra is fontosok, hiszen ezen rendszerek rengeteg állandó változóként megadott promptot tartalmaz. Ezen fix promptok szerves részét képezik az adatfeldolgozásnak és a különböző integrációknak. Amennyiben egy felhasználó interakciót kezdeményez az LLM alapú megoldásunk felé az jelen esetünkben egy prompttal indul, amely még kiegészítésre kerül a vektor adatbázisban tárolt releváns információval és más egyéb integrált információforrástól származó információval. Fontos, hogy nincs tökéletes prompt, és nincs tökéletes képlet hozzájuk, mivel ez nagyban függ a feldolgozott adattól és az egyéb tényezőktől. Érdemes az aktuális ajánlásokat átnézni és sokat tesztelni, kísérletezni a promptokkal, hogy megtaláljuk a mi igényeinknek és környezetünknek a leginkább megfelelőt.[17]
 10. **RAG támogatott LLM alapú Chatbot publikálása:** Az RAG támogatott LLM alapú chatbot-ot a megfelelő tesztelés és felhasználói oktatás után megnyitásra kerül a felhasználók részére. A tesztelés során vagy akár már a publikálás után érkező felhasználói visszajelzések alapján finomhangolást iteratív jelleggel rendszeresen el kell végezni. [18]
 11. **RAG támogatott LLM alapú Chatbot integrálása a munkafolyamatokba:** Amennyiben elfogadható válaszokat ad a chatbot, akkor átfogóan át kell tekinteni, hogy mely területeken érdemes először aktívan integrálni. Itt érdemes figyelembe venni azt, hogy mely csapatoknál használható leginkább a megoldás és hogy mely csapatoktól érkehetnek hasznos visszajelzések, ideértve azokat, akik valamilyen módon találkoztak már információbiztonsági követelményekkel, vagy kellett volna már találkozniuk. Az információbiztonsági szakemberek különösen nagy segítség tudnak lenni, hiszen ők a tapasztalatukból adódóan releváns visszajelzést tudnak adni. A felmérés és a bevezetés során fontos a vezetői támogatás és érdemes lehet projekt menedzserek segítségével végezni ezen bevezetést. Fontos felkészíteni a felhasználókat a rendszer esetleges hibáira és informálni őket a problémák bejelentésére és visszajelzésére szolgáló kommunikációs csatornákról és annak használatáról. [12], [13], [18]

12. **Folyamatos naprakészen tartás, folyamatos fejlesztés:** Mivel a szervezettel vagy részleggel szemben érvényes követelmények szignifikánsan változhatnak, ezeknek dokumentáció oldalról is változásként kell megjeleníteniük, amelynek pedig a vektor adatbázisban történő változásnak kell lekövetnie. A rendszerből adódóan folyamatosan változni fog a mögötte lévő dokumentum rendszer, így folyamatos finomhangolásra és fejlesztésre van szükség a rendszer esetén. Ezen változások kezelésére megfelelő folyamatokat kell létrehozni. [12], [13], [18]

AZ RAG TÁMOGATOTT LLM-ALAPÚ MEGFELELŐSÉG ELŐNYEI ÉS KOCKÁZATAI

Ezen RAG támogatott LLM megoldás működéséből adódóan rengeteg szervezeti információval dolgozik, a publikációban ennek kockázatait nem tárom fel, hiszen nagyon eltérő módon állnak a szervezetek a nagyon eltérő szinten bíznak meg a LLM modellekben és a mögöttük álló cégekben adatbiztonsági szempontból. A bemutatott megoldás kivitelezhető a szervezet által üzemeltetett AI modellek segítségével is, de ennek komplexitása miatt ennek részleteit nem taglalja a publikáció. [12], [13], [18]

Lehetséges előnyök

Az RAG támogatott LLM a különböző IT rendszereket érintő folyamatokba történő integrálása számos előnyt kínál :

- **Gyorsabb információbiztonsági folyamatok:** Az LLM-ek azonnali hozzáférést biztosítanak az információbiztonsági követelményeket érintő információkhoz, csökkentve a manuális kutatásra vagy más a követelményeket ismerő emberekre várással töltött időt, ezzel felgyorsítva és leegyszerűsítve a követelmények tisztázását.[7], [9], [19]
- **Konzisztens naprakész információ forrás:** Az egységes és felügyelt információs háttér miatt RAG támogatott LLM-ek közel valós időben tudnak reagálni a dokumentált követelményekben történő változásokra és tudják biztosítani az információbiztonsági követelmények és kontrollok konzisztens, egységes alkalmazhatóságát különböző projektekben, minimalizálva az információ terjedéséből adódó késések okozta eltéréseket. [7], [9], [19]
- **Csökkenő emberi erőforrás igény:** Az automatizálás miatt csökken az információbiztonsági szakemberek által manuálisan elvégzett munka, csökken a hozzájuk érkező információbiztonsági szabályzattal kapcsolatos kérdések száma vagy komplexitása.
- **Könnyű hozzáférhetőség:** Az RAG támogatott LLM-ek könnyen hozzáférhetőek, időzóna vagy munkaidő korlátok nélkül. Lehetővé téve a 24/7 történő kérdések feltevélet, így a felhasználók számára is sokkal kézenfekvőbb és kényelmesebb lesz ezt használni, ezzel csökkenthető az információbiztonsági szakemberek felé történő kitettség. [7], [9], [19]

Lehetséges kockázatok

Az előnyök ellenére az RAG támogatott LLM-ekre való túlzott támaszkodás kockázatokkal járhat, amely szervezetenként más súllyal jelentkezik, és szervezetenként eltérő módon lehet ezekre reagálni:[20]

- **Lehetséges félreértelmességek:** Az RAG támogatott LLM-ek megfelelő ellenőrzési visszacsatolási mechanizmusok nélkül vagy nem megfelelően széleskörűen dokumentált környezet esetén félreértelmezhetik az összetett megfelelőségi kontextusokat vagy a felhasználó által használt szakkifejezéseket. [21]
- **Pontatlanságok bonyolult helyzetekben:** Bonyolult információbiztonsági követelményekkel kapcsolatos szituációkban az LLM-ek visszacsatolási mechanizmusok nélkül egyszerűsített vagy hibás tanácsokat adhatnak, ami szakértői beavatkozást igényelhet.[6], [21]
- **A kontextuális árnyalatok hiánya:** Az LLM-ek nem mindig képesek teljes mértékben megérteni a szervezeti sajátosságokat, amennyiben nem megfelelően vannak azok dokumentálva, így előfordulhat, hogy az ajánlásaik nem teljesen illeszkednek az adott vállalat valóságához. [6], [21]

Javaslatok a kockázatok csökkentésére

Az RAG támogatott LLM-ek pontosságával és részletességével kapcsolatban szervezetenként eltérő az igény. Egy alap RAG támogatott LLM chatbot összeállítása publikus felhőben elérhető menedzselt komponensekből igényel nagy emberi erőforrást, viszont a rendszer finomhangolása és a különböző ellenőrző folyamat megtervezése és egyéb kiegészítő információforrás integrálása már több erőforrást igényel. Fontos megtalálni a szervezet méretének és igényeinek megfelelő optimális állapotot. Az RAG támogatott LLM-ek használatának optimalizálása és a kapcsolódó kockázatok mérséklése érdekében az alábbi ajánlásokat érdemes követni: [6], [21]

- **Felhasználók oktatása, felhasználói tananyag:** Az oktatás segítségével elősegíthető, hogy a végfelhasználók optimálisabban tudják használni a rendszert és annak előnyeit. Kiemelten fontos, hogy a végfelhasználók megfelelő kritikával állnak a rendszer válaszaival szemben. Ezen felül fontos, hogy a felhasználók megfelelő promptokat használjanak, így arra érdemes külön kitérni az oktatások során. Megfelelő használati utasítás létrehozása és közzététele szintén nagy segítség az átlag felhasználónak.[19], [21]
- **LLM alapú ellenőrzési visszacsatolás:** Különböző pontjain a rendszernek visszacsatolási pontokat hozhatunk létre, amely a bemenetet, az integrált forrásokat és a kimenetet is validálhatja.
- **Automatizált válasz minőség ellenőrzés:** LLM megoldások és egyéb adattudományi megoldások segítségével lehet elemezni a válaszokat és azoknak minőségét a kérdés függvényében.
- **Folyamatos visszacsatolási folyamat:** Folyamat létrehozása, amely hatékonyan kezeli végfelhasználók és szakértők visszajelzésének feljegyzését és azok alapján hatékonyan kezelhető és menedzselhető finomhangolások létrehozását. [22]
- **Szakértői felülvizsgálat:** Az információbiztonsági szakértők ellenőrizhetik véletlenszerűen proaktív módon vagy reaktív módon a felhasználói visszajelzésekre vagy az automatizált minőségellenőrző folyamatok kimenetelei alapján.[22]

Metrikák a hatékonyság értékelésére

Szervezetenként és integrációs típusonként eltérő, hogy mely metrika milyen fontossággal jelenik meg a szervezet számára, főbb metrikák a teljesség igénye nélkül:

- **Válaszidő:** Az általános felhasználói elégedettség miatt is kiemelten fontos a válaszidő, hiszen amennyiben túl lassú a rendszer, akkor az felhasználói frusztrációt és ellenállást okozhat. Ezen felül amennyiben más rendszerekkel van integrálva a rendszer hibákat és egyéb problémákat okozhat a túl nagy válaszidő.[23], [24]
- **Relevancia:** Az LLM kimeneteinek a kérdéssel szemben való relevancia értékelése. Ez az ellenőrzés történhet automatizált módon egy másik LLM segítségével és ezen felül a felhasználó által. [21], [23], [24]
- **Hibaarányok:** Az LLM által generált tartalomban előfordulhatnak hibák, itt figyelhetjük a hibák gyakoriságát és típusait. Ez lehet túl rövid válasz, vagy hibaüzenet is, például egy szolgáltatás kiesés miatt vagy egyes komponensek hibája miatt. [21], [23], [24]
- **Követelménnyel szemben való illeszkedés:** LLM kimenetének vizsgálata, hogy mennyire felel meg a releváns megfelelőségi szabványoknak és előírásoknak. Ez az ellenőrzés történhet automatizált módon egy másik LLM segítségével és ezen felül a felhasználó által. [21], [23], [24]

Finomhangolási lehetőségek

Az RAG támogatott LLM-ek rengeteg mozgó paraméterrel és elemmel rendelkeznek, amiknek finomhangolása szükséges a megfelelő működés elérése érdekében, ez közé tartoznak az alábbi pontok:

- **Dokumentumok optimalizálása:** Amennyiben nem megfelelő szervezeti dokumentáció alapján dolgozunk, az problémát okozhat, így fontos a megfelelő lefedettség a szervezet valóságának. Amennyiben átfednek a dokumentumainkban lefedett információk, de ezek szükségesek, akkor ezen dokumentumok között prioritási sorrendet érdemes megadni. [21], [23], [24]
- **Egyéb integrációk optimalizálása:** Meglévő integrációk felülvizsgálata és esetleges egyéb integrációs források felderítése, amennyiben még több kontextusra van szükségünk.[18]
- **Fix prompt finomhangolás:** A rendszernek szerves része a különböző fix promptok, így ezek finomhangolása és módosítása több iterációt igényelhet, ez nagy mértékben változhat amennyiben más paraméterek változnak. [25], [26]
- **Paraméter beállítás:** A rendszer több pontján is különböző paraméterek használ, az LLM kimeneteinek megfelelő befolyásolása érdekében. Ezen paraméterek lehetnek például hosszúságot, összetettséget, kreativitást, formátumot és hangnemet meghatározó elemek, hogy biztosítsuk a következetességet és az illeszkedést az elvárásainkkal szemben.[23], [25]

KÖVETKEZTETÉS

Az RAG támogatott nagy nyelvi modellek (LLM-ek) bevezetése és alkalmazása az információbiztonsági követelmények feldolgozásában jelentős előrelépést jelenthet az IT-rendszerek tervezési és kivitelezési folyamataiban. Az LLM-ek képességei lehetővé teszik a gyors és konzisztens információhozzáférést, amely kulcsfontosságú a komplex megfelelőségi követelmények betartásához. Ezek az eszközök nemcsak a tervezési ciklusok gyors

sításában játszanak szerepet, hanem az egységes információforrás biztosításával a követelmények pontosabb értelmezését és alkalmazását is elősegítik. Az azonnali válaszok révén az LLM-ek hozzájárulnak a tervezési és fejlesztési folyamat hatékonyságának növeléséhez, ugyanakkor csökkentik a mulasztásokból eredő hibák valószínűségét.

Mindezek ellenére az RAG támogatott LLM-ek használata nem mentes a kockázatoktól. Az összetett információbiztonsági kontextusok félreértelmezésének lehetősége, a pontatlan válaszok kockázata és a szervezeti sajátosságok nem megfelelő figyelembevétele mind olyan tényezők, amelyek komoly kihívást jelenthetnek a rendszer sikeres alkalmazásában. Az ilyen típusú rendszerek akkor működnek hatékonyan, ha megfelelő visszacsatolási mechanizmusokkal, szakértői felügyelettel és folyamatos finomhangolással egészítik ki azokat.

A dokumentumok rendszerezett kezelése és naprakészen tartása, az LLM-ek paramétereinek folyamatos optimalizálása, valamint a felhasználók megfelelő oktatása elengedhetetlen elemei a sikeres implementációnak. Az LLM-ek alapvető képességei a releváns tartalom generálására és az információbiztonsági követelmények értelmezésére számos gyakorlati előnnyel járnak, azonban az emberi kontroll és kritikus gondolkodás továbbra is nélkülözhetetlen marad.

A kutatás során bemutatott megközelítés világosan rávilágít arra, hogy az LLM-ek alkalmazása nem csupán a megfelelőségi követelmények feldolgozását, hanem azok szervezeti környezethez való igazítását is jelentősen megkönnyítheti. Az olyan technológiák, mint az RAG támogatott LLM-ek, lehetőséget teremtenek az információs struktúrák hatékonyabb kezelésére, miközben csökkentik a manuális erőforrásigényt és a hibák valószínűségét. Az RAG támogatott LLM-ek szerepe az információbiztonsági megfelelőségkezelésben folyamatosan bővülni fog, ahogyan ezek az eszközök tovább fejlődnek és még elérhetőbbek lesznek. Az automatizált megoldások és az emberi szakértelem együttes alkalmazásával a szervezetek képesek lesznek hatékonyabb biztonsági szintet fenntartani, miközben hatékonyabbá teszik az informatikai rendszerek tervezését, üzemeltetését érintő folyamataikat. A technológia adta lehetőségek kiaknázása ugyanakkor csak akkor lesz sikeres, ha azt átgondolt és alaposan megtervezett folyamatokkal támogatják, amelyek figyelembe veszik a szervezet specifikus igényeit és kihívásait.

FELHASZNÁLT IRODALOM

- [1] B. Kenyon, 'Nine steps to success: an ISO 27001: 2022 implementation overview', 2024.
- [2] P. Nurmi, 'Enhancing ISO 27001: 2022 Implementation Through Project Management', 2024.
- [3] Y. Kurii and I. Oprisky, 'Analysis and Comparison of the NIST SP 800-53 and ISO/IEC 27001: 2013', *NIST Spec Publ*, vol. 800, no. 53, p. 10, 2022.
- [4] R. Islam and I. Ahmed, 'Gemini-the most powerful LLM: Myth or Truth', presented at the 2024 5th Information Communication Technologies Conference (ICTC), IEEE, 2024, pp. 303–308.
- [5] B. Chand and B. Patel, 'RAG-based Chat Application using LLMs: A Case Study of Vikram Sarabhai Library IIM Ahmedabad', 2024.

- [6] ‘What is Retrieval-Augmented Generation (RAG)?’, Google Cloud. Accessed: Nov. 06, 2024. [Online]. Available: <https://cloud.google.com/use-cases/retrieval-augmented-generation>
- [7] ‘RAG Engine overview | Generative AI on Vertex AI’, Google Cloud. Accessed: Nov. 06, 2024. [Online]. Available: <https://cloud.google.com/vertex-ai/generative-ai/docs/rag-overview>
- [8] ‘Vertex AI with Gemini 1.5 Pro and Gemini 1.5 Flash’, Google Cloud. Accessed: Nov. 06, 2024. [Online]. Available: <https://cloud.google.com/vertex-ai>
- [9] S. B. Islam, M. A. Rahman, K. Hossain, E. Hoque, S. Joty, and M. R. Parvez, ‘OPEN-RAG: Enhanced Retrieval-Augmented Reasoning with Open-Source Large Language Models’, *ArXiv Prepr. ArXiv241001782*, 2024.
- [10] R. Sahu and M. Speretta, ‘Statistical Word Analysis to support the Semiautomatic Implementation of the NIST 800-53 Cybersecurity Framework’, 2024.
- [11] ‘What is ISO 27001? – TechTarget Definition’. Accessed: May 20, 2023. [Online]. Available: <https://www.techtarget.com/whatis/definition/ISO-27001>
- [12] D. Moskowitz and D. Nichols, *A Practitioner’s Guide to Adapting the NIST Cybersecurity Framework: Create, Protect, and Deliver Digital Business Value Series*. Tso, the Stationery Office, 2023. [Online]. Available: <https://books.google.hu/books?id=wd2EzweECAAJ>
- [13] C. Bartsch, *Information Security Based on ISO 27001 Strategies: A Leadership Introduction to Information Security*. Amazon Digital Services LLC - Kdp, 2023. [Online]. Available: <https://books.google.hu/books?id=0shi0AEACAAJ>
- [14] ‘What is RAG? - Retrieval-Augmented Generation AI Explained - AWS’, Amazon Web Services, Inc. Accessed: Nov. 06, 2024. [Online]. Available: <https://aws.amazon.com/what-is/retrieval-augmented-generation/>
- [15] HeidiSteen, ‘RAG and generative AI - Azure AI Search’. Accessed: Nov. 06, 2024. [Online]. Available: <https://learn.microsoft.com/en-us/azure/search/retrieval-augmented-generation-overview>
- [16] ‘RAG with databases on Google Cloud’, Google Cloud Blog. Accessed: Nov. 06, 2024. [Online]. Available: <https://cloud.google.com/blog/products/ai-machine-learning/rag-with-databases-on-google-cloud>
- [17] G. Team *et al.*, ‘Gemini 1.5: Unlocking multimodal understanding across millions of tokens of context’, *ArXiv Prepr. ArXiv240305530*, 2024.
- [18] ‘Build a RAG chatbot with GKE and Cloud Storage | Kubernetes Engine’, Google Cloud. Accessed: Nov. 06, 2024. [Online]. Available: <https://cloud.google.com/kubernetes-engine/docs/tutorials/build-rag-chatbot>
- [19] M. Kulkarni, P. Tangarajan, K. Kim, and A. Trivedi, ‘Reinforcement Learning for Optimizing RAG for Domain Chatbots’, *ArXiv Prepr. ArXiv240106800*, 2024.
- [20] B. Jonkhout, ‘Evaluating Large Language Models for Automated Cyber Security Analysis Processes’, 2024.
- [21] Z. Ji, T. Yu, Y. Xu, N. Lee, E. Ishii, and P. Fung, ‘Towards mitigating LLM hallucination via self reflection’, presented at the Findings of the Association for Computational Linguistics: EMNLP 2023, 2023, pp. 1827–1843.

- [22] J. Wei, Y. Yao, J.-F. Ton, H. Guo, A. Estornell, and Y. Liu, ‘Measuring and reducing llm hallucination without gold-standard answers via expertise-weighting’, *ArXiv Prepr. ArXiv240210412*, 2024.
- [23] D. Banerjee, P. Singh, A. Avadhanam, and S. Srivastava, ‘Benchmarking LLM powered chatbots: methods and metrics’, *ArXiv Prepr. ArXiv230804624*, 2023.
- [24] T. Hu and X.-H. Zhou, ‘Unveiling LLM Evaluation Focused on Metrics: Challenges and Solutions’, *ArXiv Prepr. ArXiv240409135*, 2024.
- [25] L. Li, Y. Zhang, and L. Chen, ‘Prompt distillation for efficient llm-based recommendation’, presented at the Proceedings of the 32nd ACM International Conference on Information and Knowledge Management, 2023, pp. 1348–1357.
- [26] J. Zamfirescu-Pereira, R. Y. Wong, B. Hartmann, and Q. Yang, ‘Why Johnny can’t prompt: how non-AI experts try (and fail) to design LLM prompts’, presented at the Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems, 2023, pp. 1–21.