

SZAKDOLGOZAT

BUDAPESTI GAZDASÁGI EGYETEM
KÜLKERESKEDELMI KAR
Nemzetközi Gazdálkodás Szak
NAPPALI munkarend
NEMZETKÖZI ÜZLETFEJLESZTÉS specializáció

AI-alapú megoldások a pénzügyi kockázatelemzésben: potenciál és etikai kihívások

Belső konzulens: Hrabovszki Ágnes Zsuzsa

Készítette: Virág Evelin

Budapest, 2025

NYILATKOZAT

AlulírottVirág Evelyn..... büntetőjogi felelősségem tudatában nyilatkozom, hogy a szakdolgozatomban foglalt tények és adatok a valóságnak megfelelnek, és az abban leírtak a saját, önálló munkám eredményei.

A szakdolgozatban felhasznált adatokat a szerzői jogvédelem figyelembevételével alkalmaztam.

Ezen szakdolgozat semmilyen része nem került felhasználásra korábban oktatási intézmény más képzésén diplomaszerzés során.

Tudomásul veszem, hogy a szakdolgozatomat az intézmény plágiumellenőrzésnek veti alá.

Budapest, 2025. év aug...... hónap5..... nap

.....Virág Evelyn.....

hallgató aláírása

Tartalomjegyzék

1.	Bevezetés	6
1.1.	Megfogalmazott hipotézisek.....	7
2.	Szakirodalmi áttekintés.....	7
3.	Alkalmazott kutatási módszerek.....	11
4.	Mesterséges intelligencia alkalmazása a kockázatelemzésben	12
4.1.	Mesterséges intelligencia keretrendszere és használati esetei	12
4.2.	Az AI felhasználási területei a kockázatelemzésben	15
4.3.	A mesterséges intelligencia előnyei és hátrányai a kockázatelemzésben	18
5.	Adatvédelem és adatbiztonság	19
5.1.	Személyes adatok védelme.....	19
5.2.	Társadalmi hatások	20
6.	A kérdőív elemzése	23
6.1.	Demográfiai adatok	23
6.2.	Kockázatok.....	23
6.3.	Mesterséges intelligencia	25
6.4.	AI-alapú megoldások hatása a kockázatelemzésre.....	26
6.5.	Etikai kihívások és adatvédelem.....	27
6.6.	Jövőbeli kilátások.....	29
6.7.	Összegzés.....	30
7.	Esettanulmányok.....	30
7.1.	Sikeres mesterséges intelligencia alapú kockázatelemzési projektek	31
7.1.1.	Shawmut Design and Construction.....	31
7.1.2.	RAZE Finance	32
7.1.3.	Upstart.....	33
7.2.	Kudarok és tanulságok.....	35
7.2.1.	Diszkriminatív AI rendszerek és azok következményei	35
7.2.2.	AI projektek kudarcaránya.....	36
8.	Összefoglalás.....	36

9. Irodalomjegyzék.....	39
10. Ábrajegyzék.....	42
11. Mellékletek.....	42

1. Bevezetés

A mesterséges intelligencia az egyik legdinamikusabban fejlődő terület a technológia világában, amely számos iparágat forradalmasított napjainkban. Különös áttörést ért el a mesterséges intelligencia a kockázatelemzés területén, tekintve, hogy képes hatalmas adatbázisokból dolgozva gyorsan eredményekhez jutni az elsőre nem szembetűnő összefüggések azonosításával is.

A hagyományos kockázatelemzési módszerek már nem feltétlenül tudnak lépést tartani a manapság jellemző gyors változásokkal, valamint a növekvő adatmennyiség is problémát okozhat az optimális működésben. Ezzel szemben a mesterséges intelligencia rendszerek a prediktív elemzés és a gépi tanulás segítségével dinamikusan képesek alkalmazkodni, képesek rettentő gyorsan reagálni a változásokra. A pénzügyi szektorban kifejezetten a hitelkockázat értékelés, a csalásfelderítés esetében nyújthat hatalmas támogatást, azonban az építőiparban és a bűnügyi szférában is segítségül hívhatjuk, ezeken a területeken már rendelkezünk mérhető adatokkal is.

Minden előny ellenére azonban új etikai és társadalmi kérdést is magával hordoz a mesterséges intelligencia alkalmazása. A kockázatelemző rendszerek gyakran a fekete doboz jelenséggel találkoznak, amelynek a döntései sokszor átláthatatlanok. Ez különösen akkor óriási probléma, ha egy ilyen okból kifolyólag nem jut valaki hitelhez, vagy potenciális elkövetőként azonosít valakit egy rendészeti rendszer. Az ilyen döntések mögötti adatok torzítottak lehetnek, amelyek komoly etikai dilemmákhoz vezethetnek.

A mesterséges intelligencia működése tehát kétoldalú: a kockázatelemzési folyamatokat hatékonyabbá, gyorsabbá teszi, de eközben etikai, társadalmi és esetenként még jogi kérdéseket is felvethet. A dolgozatom célja, hogy bemutassa a mesterséges intelligencia alapú kockázatelemzési rendszerek működését és lehetőségeit, miközben részletesen megvizsgálja az ezekhez kapcsolódó etikai kérdéseket is. A dolgozat további célja, hogy feltérképezze a mesterséges intelligencia alkalmazásának jövőbeli lehetőségeit, valamint, hogy milyen szempontokat lehet érdemes figyelembe venni a társadalmi igazságosság és biztonság biztosítása érdekében.

Dolgozatomban esettanulmányokat, szakirodalmat, valamint kérdőíves felmérést is segítségül hívok, hogy minél átfogóbb, árnyaltabb képet kaphassunk a mesterséges intelligencia szerepéről a kockázatelemzésben. A kutatás célja, hogy feltárja az etikusabb mesterséges

intelligencia használatot, hogy ne csak technikai, de társadalmi szempontból is elfogadható legyen az ilyen rendszerek alkalmazása.

1.1. Megfogalmazott hipotézisek

- 1) Első felvetésem arra irányul, hogy az AI alapú kockázatelemző rendszerek pontosabb előrejelzéseket, analitikákat képesek nyújtani, mint a hagyományos statisztikai modellek.
- 2) Második hipotézisem, hogy a mesterséges intelligencia használata átláthatatlanságot eredményez, így csökkenti a bankokkal szembeni bizalmat.
- 3) Harmadik felvetésem már az adatvédelemre vonatkozik. Tekintve, hogy a pénzügyi kockázatelemző rendszerek gyakran személyes adatokat dolgoznak fel etikai szempontból jelentős kihívásokat vethet fel a mesterséges intelligencia alkalmazása.

A hipotéziseim azért relevánsak, mert elemzésükkel a jövő lehetőségeit vizsgálom a pénzügyi kockázatelemzésben. Ezekkel a felvetésekkel nemcsak a technológiai hatékonyságra, pontosságra térek ki, számba veszem az etikai, társadalmi bizalmat érintő kérdéseket is. Hipotéziseim azért is kiemelkedően fontosak, mert a mesterséges intelligencia már a mindennapjaink részévé vált, azonban nem feltétlenül alkalmazható minden területen, vagy alapos utómunkát, átnévezést igényelnek az általa elvégzett feladatok. Összességében tehát mindegyik hipotézis olyan kulcskérdéseket helyez a fókuszba, amelyek nélkülözhetetlenek a mesterséges intelligencia felelősségteljes, biztonságos alkalmazásához.

2. Szakirodalmi áttekintés

A fejezet célja a dolgozat tudományos kontextusba történő helyezése, így ennek értelmében a továbbiakban a releváns szakirodalom segítségével ismertetem azokat az alapfogalmakat, amelyek a dolgozat későbbi részeiben megjelennek.

Mindenekelőtt a kockázatot és a mesterséges intelligenciát definiálom.

A **kockázat** szóra rengeteg definíciót találhatunk, mind az interneten, mind a könyvekben. Ha a hétköznapi életünket vesszük alapul, a kockázat kifejezés elsősorban negatív jelentésárnyalattal bír, azonban a kockázatok valójában objektívan mérhetőek és segítségével elkerülhetjük a negatív hatásokat, akár még csökkenthetjük is a kockázat nagyságát (Erdős & Mérő, 2010) A pénzügyi intézmények számára különösen fontos a

kockázatokat megfelelően kezelünk és csoportosítanunk. A kockázatokat többféle módon, különböző szempontok alapján azonosíthatjuk. A klasszikus csoportosítás szerint két nagy csoportra oszthatóak fel: pénzügyi (mérhető) kockázatok, illetve üzleti (nem mérhető) kockázatok (Horváth S., Koltai, & Nádasdy, 2019). A pénzügyi kockázatok közé sorolhatjuk a hitelkockázatot, a piaci kockázatot és a likviditási kockázatot, míg az üzleti kockázatok esetében beszélhetünk működési, stratégiai és reputációs kockázatról is.

A **mesterséges intelligencia** egy olyan specialitás a számítástechnikán belül, amely olyan rendszerek létrehozásával foglalkozik, amelyek képesek megismételni az emberi intelligenciát és a problémamegoldó képességeket. Ezt számtalan mennyiségű adat befogadásával és feldolgozásával teszik, illetve az emberek tanulás módját igyekeznek elsajátítani, mint például, hogy önállóan hajtsanak végre feladatokat, a felemerülő hibákat, pontatlanságokat javítsák a tapasztalatok és több adatnak való kitettség segítségével. (IBM, 2021)

Mint ahogyan már a bevezető fejezetben is kifejtettem a pénzügyi szektornak jelenleg számos lehetőséggel és kihívással kell szembenéznie: pénzügyi kockázatok, operatív kockázatok, gépi tanulás, természetes nyelvfeldolgozás, prediktív analitika, GDPR szabályozások, feketedoboz jelenség. Az alábbiakban ezeket a szakkifejezéseket és jelenségeket definiálom.

A **pénzügyi kockázatok** esetében három típust szeretnék kiemelni, amelyeket a legfontosabbnak találtam: hitelkockázat, piaci kockázat, likviditási kockázat.

A bankok által vállalt legnagyobb kockázat a **hitelkockázat**, mivel a pénzintézet a saját visszafizetési kockázatává alakította azt a kockázatot, hogy az általa meghitelezett ügyfelek a hitelt és annak kamatait nem tudják időben visszafizetni, vagy egyáltalán nem fizetik vissza. (Erdős & Mérő, 2010) A **piaci kockázatok** esetében a bankok nem tudják közvetlenül befolyásolni az árakat, hiszen ezek a részvényárfolyamok, kamatlábak, devizaárfolyamok, vagy a tőzsdén kereskedett árucikkek ingadozásából adódnak, ezáltal a pénzintézetek egyes portfólióelemeit (amelyekkel kifejezetten az a cél, hogy rövid távon profitot realizáljanak) piaci kockázatnak kitett kereskedési portfóliónak nevezzük. Ezek az elemek lehetnek kötvények, részvények, az ugyancsak idesorolható devizapozíciókkal, amelyek a devizaárfolyam változása miatti devizaárfolyam-kockázatnak vannak kitéve. (Erdős & Mérő, 2010) **Likviditási kockázat** alatt azt a jövedelmezőségi és

tőkekövetésért, ami abból származik, hogy a bank az esedékes fizetési kötelezettségeit nem tudja jelentős veszteségek nélkül teljesíteni. A banki likviditás tehát likvid eszközök, illetve pénzeszközök tartását jelenti, azonban a likvid eszközök vagy nem termelnek jövedelmet (például készpénz) vagy kevésbé kifizetődőek (például rövid bankközi kihelyezések, állampapír-fedezettű repók). (Erdős & Mérő, 2010)

Az **operatív kockázatokkal** folytatva szintén három fő csoportot emelnék ki: működési kockázat, stratégiai kockázat, reputációs kockázat. A **működési kockázat** valójában egy gyűjtőfogalom, amely magába foglalja a banki folyamatok, az emberek, vagy a rendszerek nem megfelelő működését, vagyis olyan kockázatokat sorolunk ide, amelyek akadályozzák a bank zavartalan működését, ezzel veszteséget okozva a pénzintézetnek. Az általános definíció szerint hét eseménytípust különböztethetünk meg: belső csalás, külső csalás, munkáltatói gyakorlat és munkahelyi biztonság, ügyfelek/termékek és üzleti gyakorlat, tárgyi eszközök károsodása, üzletmeneti hibák és rendszerhibák, végrehajtási, teljesítési és folyamatkezelési hibák. (Erdős & Mérő, 2010) A **stratégiai kockázatok** alapvetően a politikai vagy a gazdasági környezet megváltozásának eredményeképpen születnek meg, így jogosan gondolhatjuk, hogy nehezebb ellene védekezni. A hosszú távú célok védelme érdekében azonban azonosíthatjuk a külső kockázati tényezőket (például: piaci változások, versenynyomás), ezek segítségével módosíthatjuk az üzleti stratégiát is. Segítségünkre lehet továbbá a SWOT-elemzés is, ugyanis ezzel az elemzési módszerrel mind a szervezeten belüli adottságokat (erősségek, gyengeségek), mind a külső körülményeket (lehetőségek, veszélyek) is azonosíthatjuk. (Reketye, Törőcsik, & Hetesi, 2015) A **reputációs kockázat** abból fakadhat, hogy negatív vélemények terjengenek az érintett felektől üzlettel kapcsolatban, így csökkenhet az ügyfélkör, akár még költséges pereskedésbe is keveredhet az adott vállalat. Tekintve, hogy az érintettek elvárásai folyamatosan változnak, így a hírnévkockázat dinamikus, földrajzi régióként, csoportonként és egyénenként változó. (Risk Officer, dátum nélk.)

A **gépi tanulás** (machine learning) során a számítógépek matematikai alapmodelleket tanulnak be közvetlen utasítás nélkül. Használatukkal azonosíthatóak minták az adatokban, amelyek segítségével létrehozhatunk egy alkalmas adatmodellt az előrejelzések elkészítéséhez, analíziséhez. Valójában a gépi tanulást az AI egy

részhalmozásának tartják, legfontosabb előnyei a következők: ismeretek feltárása, adatintegrálás javítása, kockázatcsökkentés. (Microsoft, dátum nélk.)

Természetes nyelvfeldolgozás alatt a szavak és a beszéd számítógép általi feldolgozását értjük. Egyesíti a nyelvészetet és a számítástechnikát, ezeket a technikákat a mesterséges intelligenciával együtt alkalmazzák a chatbotok, vagy a digitális asszisztensek létrehozására is, mint például a Google Assistant. Ahhoz, hogy a számítógépek értelmezni tudják az emberi nyelvet azokat olyan formátumúvá kell alakítani, amelyet a számítógép is képes kezelni. A természetes nyelvfeldolgozás tehát különböző algoritmusok alkalmazását jelenti, amelyek képesek strukturálatlan adatok felvételére, majd azok strukturálttá alakítására. (Nelson, 2024)

A **prediktív analitika** rendkívüli pontossággal képes előrejelzéseket készíteni a jövőbeli eseményekről, viselkedésről. Az efféle elemzés statisztikai elemzési módszereket alkalmaz, ennek négy nagy csoportját tartjuk számon: lineáris és logisztikus regresszió, döntési fa, neurális hálózatok, illetve idősoranalízis. (Dyntell, dátum nélk.)

A **GDPR** (általános adatvédelmi rendelet) **szabályozás** az EU által elkészített és elfogadott legszigorúbb adatvédelmi és biztonsági törvény a világon. Mindaddig kötelezettségeket ró a szervezetekre, ameddig az Európai Unióban élő emberekkel kapcsolatos adatokat gyűjtenek. Amennyiben ezen szabályozást megsértik, vagy megszegik szigorú pénzbírsággal sújtja az elkövetőket a rendelet, amely akár több tízmillió eurós büntetést is jelenthet. (Wolford, dátum nélk.)

A **feketedoboz jelenség** alatt azt értjük, hogy a mesterséges intelligencia modellek anélkül döntenek, hogy magyarázatot adnának arra, hogy hogyan jutottak hozzá az adott információhoz. Az AI technológia fejlődésével két fő típus jelent meg: a fekete doboz AI és a fehér doboz (magyarázható) AI. Egyszerűen fogalmazva, az olyan mesterséges intelligencia rendszereket, amelyek belső működése, döntéshozatali munkafolyamatai és hozzájáruló tényezők nem láthatók vagy ismeretlenek maradnak az emberi felhasználók számára, fekete doboz AI-rendszereknek nevezzük. (Awati & Yasar, 2024)

3. Alkalmazott kutatási módszerek

Annak érdekében, hogy alaposan feldolgozzam a témát kvalitatív és kvantitatív kutatást is végeztem. Szakdolgozatom első részében mind nemzetközi, mind hazai szakirodalmakat, szakmai cikkeket és publikációkat, illetve jogi szabályozásokat elemeztem. A célom az volt, hogy átfogó képet kapjak a mesterséges intelligencia kínálta lehetőségekről a pénzügyi kockázatelemzésben. A cikkek kiválasztása során figyelembe vettem a cikkek publikálásának időpontját, valamint relevanciájukat a kutatási kérdéseimhez. Az empirikus kutatás megalapozásához az efféle szakirodalmi áttekintés biztos alapot teremtett meg a kutatásomhoz. Az online kérdőíves felmérésben elsősorban a pénzügyi, banki szektorban dolgozó emberek ismereteit igyekeztem feltérképezni a mesterséges intelligenciával való kapcsolatukról, ismereteikről, illetve azt is vizsgáltam, hogy a hazai bankokban már alkalmaznak-e bizonyos területeken mesterséges intelligenciát a mindennapi munkavégzés során. A kérdőívben lehetőséget adtam a saját vélemények, tapasztalatok kifejtésére nyitott kérdések segítségével. Az így létrejött adathalmazt leíró statisztikával dolgoztam fel, ennek segítségével azonosítani tudtam a felmerülő trendeket, általános tapasztalatokat.

A kérdőív elemzését követően az esettanulmányokra tértem rá, ahol olyan intézmények eredményeit, tapasztalatait vettem össze egymással, amelyek már alkalmaznak mesterséges intelligenciát a folyamataikban, mindennapos munkavégzésükben aktívan segítségül hívják a mesterséges intelligencia által nyújtott lehetőségeket. Ezen esettanulmányok segítségével mélyebben megismerhettem az AI technológia gyakorlatban való alkalmazását az esetleges kudarcok és sikerek megismerésével egyaránt. Az intézmények kiválasztásánál figyelembe vettem, hogy eltérő méretű, illetve eltérő technológiai fejlettségű szervezeteket válasszak egy árnyaltabb kép elérése érdekében. Különös figyelmet fordítottam arra is, hogy milyen típusú kockázatok kezelésére használják a mesterséges intelligenciát a mintául vett intézmények, valamint, hogy milyen adatforrásokat, milyen külső és belső befolyásoló tényezők befolyásolják az implementáció sikerességét, avagy kudarcát. Összességében a dolgozatomban elsősorban nem technikai elemzést végeztem, ugyanis nem hasonlítottam össze egymással a különböző modelleket. A kutatásom fő szempontja a már meglévő cikkek, tanulmányok eredményének bemutatása és ismertetése volt.

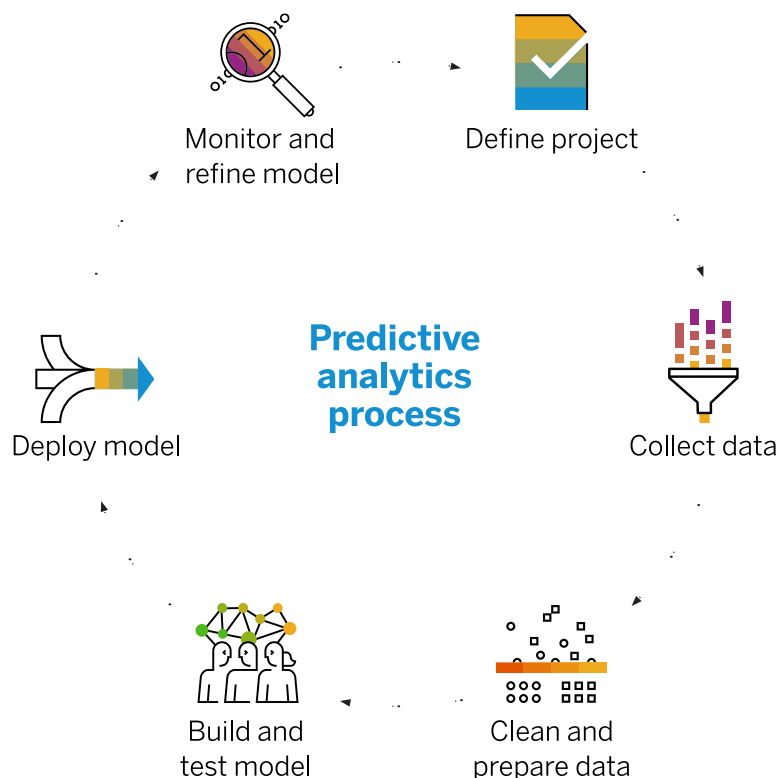
4. Mesterséges intelligencia alkalmazása a kockázatelemzésben

A mesterséges intelligencia alkalmazása szinte minden iparágban megjelent, azonban a kockázatelemzés területén kiemelkedő hatást gyakorolt az eszköz térnyerése. A hatalmas adatbázisok mintázatainak felismerése és gyors elemzése által rövidebb időt vehet igénybe a döntéshozatal, pontosabban ismerhetjük meg a kockázatok valódi mértékét. A KPMG (Magyarország egyik vezető könyvvizsgáló, adó-és üzleti tanácsadó társasága) egy 2021-es cikke szerint a hitelképesség elemzéséhez, a csalásfelderítéshez és a piaci kockázatok lehetséges előrejelzéseikhez gyakrabban alkalmazzák a mesterséges intelligenciát, feltehetően ez az elmúlt négy évben meg inkább jellemző a kockázatelemzés területén. (KPMG, 2021)

Természetesen akadhatnak problémák is a mesterséges intelligencia alkalmazása során, ezt támasztja alá a Washington Post bejegyzése is, amiben kimutatták, hogy a jelenlegi modellek gyakran hibáznak, nem teljesítenek megfelelően egyes pénzügyi feladatokban, így megkérdőjelezhetővé válik az eszköz megbízhatósága is. (Tiku, 2025) Ezeket az aggályokat tovább feszegetve a Reuters bejegyzése az etikai, jogi felelősség kérdéskörének vizsgálatát vette figyelembe, a mesterséges intelligencia modellek autonóm működése végett. (Ken, Pramode , & Weiss, 2025) Mivel a fent említett aggodalmak miatt elengedhetetlen a mesterséges intelligencia kockázatosságának vizsgálata, így a következő alfejezetekben beszeretném mutatni a mesterséges intelligencia alkalmazása mellett szóló indokokat, szeretném feltérképezni a gyakorlati felhasználási lehetőségeket, valamint fel kívánom mérni a mesterséges intelligencia esetleges buktatóit, valamint az előnyeit.

4.1. Mesterséges intelligencia keretrendszere és használati esetei

A pénzintézetek naponta több ezer adatot állítanak elő és dolgozzák fel azokat. A mesterséges intelligencia gyorsaságával az eredmények rövidebb időn belül realizálódhatnak, így a tények is korábban rendelkezésünkre állhatnak. A gyorsaság mellett azonban kulcsfontosságú az adatok precizitása is. A prediktív elemzés egy olyan analitikai módszer, amely jövőbeli eseményeket, trendeket azonosít statisztikai technikák alkalmazásával. Egy ilyen elemzés során a gépi tanulási algoritmusok, valamint az aktuális és korábbi adatok vizsgálata és elemzése (kifinomult prediktív modellezés) kerül a fókuszba. Ezt a fajta elemzési módszer elsősorban készletelőrejelzésre, erőforráskezelésre, hitelkockázati modellek kifejlesztésére használják, jelentősen csökkentheti egy vállalat műveleteit optimalizálni. (SAP, dátum nélk.)



ábra 1 Prediktív elemzés

Forrás: <https://www.sap.com/hungary/products/data-cloud/cloud-analytics/what-is-predictive-analytics.html>

A fenti ábrán a prediktív elemzés folyamatábráját láthatjuk. Első lépésként a projekt céljait kell meghatároznunk, vagyis azt, hogy mi az alapvető probléma, amire megoldást keresünk. A probléma feltárását az adatgyűjtés követi. Ezek az adatok származhatnak különböző forrásokból, akár még különböző területekről is, ebben a folyamatban óriási szerepe lehet a mesterséges intelligenciának is. Ezt követően szükséges elvégeznünk az adattisztítást, valamint az információk előkészítését az elemzésre. Ekkor van lehetőségünk pótolni a még hiányzó adatokat, valamint eltávolítanunk a nem releváns vagy idejétmúlt tájékoztatásokat is (ezzel javíthatjuk az adathalmaz minőségét). Mindezek után össze kell állítanunk a prediktív modellt és természetesen meg kell tanítanunk az adathalmazra, majd teszteket kell rajta futtatnunk a pontosság biztosítása végett. Nagyon ritkán fordul elő, hogy elsőre hibamentes modellt generálunk, ezért több iteráció is szükséges lehet. Mindezek után már üzembe is helyezhetjük a modellt, ami már készen is áll az új adatok feldolgozására. Érdeemes azonban az üzembe helyezés után rendszeresen ellenőrizni a modell teljesítményét a zökkenőmentes döntéshozatal elérése érdekében.

Az imént említett automatizálással elérhetjük, hogy csökkentsük az emberi beavatkozás szükségét. Ez leginkább a rutin vagy monitoring feladatok elvégzése során lehet fontos, mint

például a tranzakciók ellenőrzése. A KPMG álláspontja szerint a mesterséges intelligencia számos kockázatkezelési célt képes kielégíteni, a kérdés valójában nem az, hogy használják-e ezt a technológiát, inkább az, hogy hogyan. A gyors változások azt eredményezik, hogy egy vállalatnak naprakész, preventív módszerekkel és adatokkal kell rendelkezniük. Ebben óriási segítségünkre lehet a mesterséges intelligencia, ugyanis dinamikusan, aktuális adatokkal dolgozik, amely elősegíti az óriás mennyiségű adatáramlás feldolgozását. Egyre inkább elterjedt a vállalati körben a harmadik felekkel összefüggő kockázatok kezelésére vonatkozó irányítási keretrendszer (TPRM: Third-Party Risk Management), azonban a decentralizált szoftverek és szabvány nélküli gyakorlatok eredményeképp a TPRM jelentések hitelességüket veszítik. Az ilyen vállalkozások számára érdemes lehet a mesterséges intelligencia bevezetését szorgalmazni, ugyanis ezt támogatják a TPRM rendszerek, továbbá az analitikai képességek és a folyamataktualizálási lehetőségek javíthatják a jelentések hatékonyságát is. (KPMG, 2021) A gyorsaság mellett azonban elengedhetetlen a proaktivitás is, hogy a potenciális kockázatokat még azelőtt kezelhessük, hogy be is következzenek. Ezt támasztja alá a Visure Solutions blogbejegyzése is, mely szerint a gépi tanulási algoritmusok a folytonos tanulás által (múltbeli és jelen idejű adatokból) javíthatják a kockázatok előrejelzésének képességét. Vegyünk egy példát a kiberbiztonság területéről: szokatlan tevékenységek zajlanak a rendszerekben (amelyek akár még jogsértésre is utalhatnak) ezt azonosítják az algoritmusok, ezzel nagyban hozzájárulnak a gyors közbeavatkozáshoz. A példa alapján is láthatjuk, hogy az időzítés kulcsfontosságú lehet mind egy kibertámadás, működési és pénzügyi kockázat megelőzése esetében is, ezáltal ugyanis időben feltárhatjuk a lehetséges fenyegetéseket és még bekövetkezésük előtt meggátolhatjuk őket. (Visure Solutions, dátum nélk.)

Természetesen minden vállalkozásnak vannak megegyező, valamint egyedi kockázatai is, a mesterséges intelligencia használata ezért lehet ebben az esetben is kiemelkedő választás. A személyre szabott kockázatelemzéssel még pontosabb és relevánsabb adatokkal dolgozhatunk tovább, mintha egy hagyományos kockázatkezelési módszert választanánk, mint például a szegmens alapú megközelítés. Az automatizált rendszereknek köszönhetően a folyton változó szabályozások is nyomonkövethetőbbek, hiszen szinte azonnal tudnak alkalmazkodni az új kritériumoknak. Az átláthatóság érdekében a mesterséges intelligenciát alkalmazó rendszerek képesek dokumentációval alátámasztani az eredményeiket, így könnyen lekövethető, hogy egy adott döntés milyen logika mentén jött létre. Ezt nevezzük magyarázható mesterséges intelligenciának (explainable AI), amely kulcsfontosságú szerepet tölthet be a pénzügyi szektorban a megbízhatóság, a bizalom és az átláthatóság szempontjából. (PWC, dátum nélk.)

4.2. Az AI felhasználási területei a kockázatelemzésben

Az előző alfejezet segítségével alaposan megismerhettük a mesterséges intelligencia alkalmazásának előnyeivel, keretrendszerével. Ebben az alfejezetben be szeretném mutatni azokat a területeket, amelyekben már alkalmaznak mesterséges intelligenciát, még hozzá rendkívül elősegíti a folyamatokat is a pénzügyi szektorban. Első sorban a hitelképesség-elemzésben nyújt óriási segítséget a mesterséges intelligencia. A hitelkockázati profil részét képezik a különböző pénzügyi kimutatások, valamint a hitelnyilvántartások is, amelyek segítségével pontosabban és gyorsabban értékelhetjük az ügyfelek hitelkockázatát. A Magyar Nemzeti Bank állásfoglalása szerint is megbízhatóbb eredményekhez juthatunk a mesterséges intelligencia alkalmazásával a hagyományos módszerekkel ellentétben. (Dr. Szikora & Nagy , dátum nélk.) Az állandó rendelkezésre állásnak köszönhetően a rendszerek képesek valós időben azonosítani a tranzakciókat, ezzel elősegítve a potenciális csalások beazonosítását. Ezt a technológiát alkalmazza például a JP Morgan Chase a világ egyik legnagyobb bankja is a tranzakciós minták elemzésére. A rendszereik úgy lettek kialakítva, hogy részletes vásárlási profilokat készítsenek minden ügyfél számára, amelyek elősegítik észlelni a megszokott költségi magatartástól való eltérést, valamint azokat, amelyek tiltott tevékenységekre utalnak. Az egyes tranzakciók kockázati pontszámait a korábbi adatokból, illetve az említett megváltozott magatartás alapján értékelik. (Dujmovic, 2024) A JP Morgan Chase- hez hasonlóan a BlacRock, egy globális befektetés kezelő vállalat is alkalmaz mesterséges intelligenciát a befektetési portfóliójához kapcsolódó kockázatok felmérése érdekében. A szervezet Aladdin nevű platformja gépi tanulás segítségével elemzi a piaci adatokat, gazdasági mutatókat ezzel optimalizálva a befektetési döntéseket, valamint a kockázattal kiigazított hozam optimalizálását. Ez a platform a vállalat stratégiájának szerves részévé alakult, ebből a példából is láthatjuk, hogy rendkívüli szerepe lehet a mesterséges intelligenciának a pénzügyi szektorban. (Ahmad, 2024) A mesterséges intelligenciát azonban nem csak a stratégiai döntéshozatal esetén hívhatjuk segítségül, az operatív kockázatok kezelésénél is alkalmazhatjuk az esetleges rendszerhibák azonosításakor. Ehhez a természetes nyelvfeldolgozást érdemes alkalmazni, hiszen általa automatikusan elemzésre kerülnek a szabályozási dokumentumok és jelentések is. (Yallo, dátum nélk.) Szorosan érinti ezt a kérdéskört a pénzmosás ellen való küzdelem, valamint az ügyfélátvilágítás is. Pénzmosás abban az esetben merülhet fel, amikor az ügyfelek igyekeznek eltitkolni, elrejteni a jövedelmük eredetét. Ilyen esetben az elkövetők megpróbálhatják igénybe venni a bankok szolgáltatásait, hogy „átmossák”, legálissá tegyék a pénzük eredetét. A

terrorizmusfinanszírozás esetében másról beszélünk, olyankor a pénz eredete általában legális, így -a pénzmosással szemben- a cél annak felhasználása terrorfinanszírozásra, nem pedig az eredetének az eltitkolása. (OTP Bank, dátum nélk.) Az ügyfélátvilágítás, másnéven know your customer (KYC) célja, hogy alaposan megismerjük kivel üzletelünk. A folyamatban általában három lépése van. Elsőként a jogalany ellenőrzését végezzük el, ezzel megállapítjuk, hogy a jogszerű jogi személy magánszemély vagy szervezet. Ezt követően a profilalkotással frissítjük a kockázati profilt az alanyhoz kapcsolódó negatív vagy pozitív események, információk alapján. Az alany efféle szűrése segíti meghozni az üzleti döntéseket a kockázati profil, a kormányzati szankciók, valamint a tulajdonjog és irányítás segítségével. A fent említett folyamatok mind helyettesíthetőek egy mesterséges intelligencia alapú munkafolyamattal. A Moody's egy holding társaság, amely pénzügyi kutatással, elemzéssel, valamint hitelminősítéssel foglalkozik például igyekszik kialakítani egy KYC eljárást, amely a mesterséges intelligenciát tartalmazza végpontként. Az elképzelés az, hogy a felhasználók ezután egy chatben kaphatnák meg a vizsgálattal kapcsolatos információkat, eredményeket. A nehézséget az adja a kivitelezésben, hogy jelenleg kihívást jelent ilyen módon tömeges átvilágítást végezni. (Moody's, 2024) A chatbotok ötlete azonban nem a Moody's-tól származik, hiszen már számos mesterséges intelligenciát alkalmazó virtuális asszisztenst és chatbotot alkalmaznak az éjjel-nappal elérhető ügyfélszolgálat érdekében. Ezek a rendszerek képesek javítani az ügyfélélményt, hiszen gyorsan, pontosan válaszolnak és még az ügyfélszolgálati költségeket is csökkenthetik. Emberi közbeavatkozás nélkül képesek válaszolni mind tranzakciós, számlaegyenleges vagy akár még hitelezési kérdésekre is. (Crowdy.AI, dátum nélk.) A személyes adatok védelme érdekében egy rendkívül erős biztonsági protokollt alkalmaznak a pénzintézetek. Ezt elérhetik többfaktoros hitelesítéssel a számlaadatok elküldése, esetenként tranzakciók végrehajtása előtt is. Természetesen az ügyfél interakciók titkosításra kerülnek, illetve csak a releváns adatokban kereshetnek a chatbotok. (Jasińska & Puczyk, 2025)

Area	Use Case	Description
Front office	Chatbot	AI chatbot based on natural language processing (NLP)
	Intelligent marketing tool	Deep learning tool providing automated personalized offering in the app
	Real estate investment adviser	AI tool estimating real estate property prices and rentals
	NAV Planner	AI powered robo-adviser for market investment opportunities
Middle office	Credit score assessment	Free tool accessible online and in app for clients to assess their credit score
	Real-time loan	Using real time payment data and a risk management system to provide loans instantaneously
Back office	AI transaction surveillance	Real-time detection of abnormal payments and alerts users via mobile app

ábra 2 Példák a digitális bankok által bevezetett mesterséges intelligenciára

Forrás: <https://neontri.com/blog/best-banking-chatbots/>

A fenti ábra azt szemlélteti, hogy a mesterséges intelligencia egyre meghatározóbb szerephez kerül a pénzügyi szektorban a hatékonyság növelése, a költségcsökkentés, valamint az ügyfélélmény javítása érdekében. A táblázat jól bemutatja a mesterséges intelligencia egyes irodai területeken való alkalmazásának lehetőségeit különböző funkciókra és célokra egyaránt. A front office, azaz az ügyfélkapcsolati, értékesítési területek esetében a közvetlen kapcsolattartás elérése érdekében alkalmazzák. Ebben a szektorban igazán széles a mesterséges intelligencia megjelenésének skálája. A korábbiakban már kifejtettem a chatbotok adta lehetőségeket, azonban itt is fontosnak tartom megemlíteni a lehetőséget, mivel ezen a területen kiemelten javíthatja az ügyfélélményt, valamint tehermentesítheti az ügyfélszolgálatos kollégáinkat is. A chatbotokkal szemben az intelligens marketing eszközök a mély gépi tanulás (deep learning) technológiával működnek, ami a gépi tanulási módszer egy alcsoportja, amelyben neurális hálózatokat alkalmazunk. Ezek a neurális hálózatok igyekeznek az emberi biológiai idegrendszer információfeldolgozását utánozni. Ebben az esetben a neuronok, vagyis a rengeteg összefüggő elem végzik a számításokat. A neuronok rétegeiből épülnek fel a neurális hálózatok, vagyis az információ csak rétegről rétegre, egy irányba terjedhet. A mély gépi tanulás során olyan neurális hálózatról beszélünk, amelyben több rejtett réteg is található. Ezzel elérhetjük azt, hogy minden réteg különböző, egymástól egyre absztraktabb képét adja a bemeneti adatnak. (inf.u-szeged, dátum nélk.) Ez a fajta

tanulási módszer lehetővé teszi, hogy személyre szabott ajánlatokat tudjunk prezentálni az ügyfeleknek a viselkedésük, tranzakciós adataik alapján. Továbbá a célzott marketinggel növelhetjük ügyfeleink elköteleződését, és az értékesítést is hatékonyabban végezhetjük. Az ingatlanpiacon is nagy segítségünkre lehet egy mesterséges intelligencia alapú rendszer, ugyanis képes megbecsülni a reális bérleti díjakat, képes értékbecsléseket végezni, valamint a piaci trendeket is előrejelzi. Ezzel nem csak az ügyfelek döntéshozatalát segíthetjük elő, segítségével új lehetőségeket kínálhatunk a pénzintézetek számára is. Ehhez hasonló a nettó eszközérték (NAV planner) tervező rendszerek funkciója is, azonban abban az esetben nem az ingatlanárakról kapunk tájékoztatást, hanem a befektetési lehetőségekről az egyén kockázati profilja alapján. Ez a lehetőség leginkább azok számára lehet hasznos, akik kevésbé ismerik a befektetési lehetőségeiket, azonban szeretnének megfontolt döntést hozni. A middle office fő feladatai közé a kockázatelemzés, a pénzügyi ellenőrzés továbbá a különböző adatok kezelése tartozik. Ezen a területen a gyors adatfeldolgozás és döntéshozatal támogatása okozhat helyzetelőnyt, ahogy azt már a korábbiakban is kifejtettem. A hitelképesség során a hagyományos hitelmutatók mellett a mesterséges intelligencia rendszerek figyelembe vehetnek alternatív adatokat is, mint például a közösségi média aktivitási adatait vagy az online vásárlási szokásokból készült mintákat. Ezt alternatív credit scoringnak nevezzük. A nagy hitelinformációs ügynökségek, mint például a Experian Plc a hitelminősítéshez ezt a fajta alternatív hitelminősítést is figyelembe veszi. Olyan esetekben hatalmas segítséget nyújthat a minősítés elkészítésében, ha az adott ügyfél nem rendelkezik klasszikus hiteltörténettel. (SEON, dátum nélk.) A back office terület szinte alig észlelhető az ügyfelek számára, mégis fontos szerepet töltenek be a bankok zökkenőmentes működése érdekében. Ez a terület leginkább a tranzakciók felügyeletében tudja alkalmazni a mesterséges intelligencia adta lehetőségeket, leginkább a csalásmegelőzés érdekében alkalmazzák, ahogy azt már korábban is kifejtettem.

4.3. A mesterséges intelligencia előnyei és hátrányai a kockázatelemzésben

Mostanra már megbizonyosodhattunk róla, hogy a mesterséges intelligencia alkalmazása a kockázatelemzésben rendkívül sok előnnyel járhat, azonban hátrányokat és kockázatokat is hordozhat magában. Előnyei közé tartozik a hatékonyság és az automatizálás, amely csökkenti az emberi munkaerő igényét, így gyorsabban, pontosabb eredményt kaphatunk. (Visure Solutions, dátum nélk.) Mivel a mesterséges intelligencia rendszerek könnyen skálázhatóak, így a vállalatoknak lehetőségük nyílik gyorsan alkalmazkodni a piaci változásokra. Ez a hatalmas adatgyűjtés azonban felvethet etikai, adatvédelmi kérdéseket.

Ezen személyes adatok begyűjtésénél elengedhetetlen, hogy biztosítsuk a védelmüket és hogy betartsuk az etikai normákat. Az automatizálhatóság miatt azonban megszűnhetnek bizonyos munkakörök, pozíciók is, emiatt egyes munkavállalóknak szükséges lehet új készségeket elsajátítaniuk, vagy teljesen más munkát végezniük, mint korábban. (Hovsepian, 2023) Alapvetően a mesterséges intelligencia alapú rendszerek megbízhatóak, azonban, ha meghibásodik a technológia, a vállalat pedig „csak egy lábon áll”, akkor el is vesztheti a vállalkozás azt a rugalmasságot és alkalmazkodóképességet, amit ezek a rendszerek biztosítanak. (Leitner, Singh, van der Kraaij, & Zsámboki, 2024) Mivel az AI rendszerek döntéshozatalai gyakran fekete dobozként működnek, így jogosan hiányolhatjuk a döntéshozatal átláthatóságát. (Stefanoski & Meier, 2024) Ez akkor lehet különösen nagy probléma, ha a döntéseket szükséges lenne indoklással ellátni, illetve, ha ellenőrizni szeretnénk az adatok hitelességét. Ilyen esetben jogosan merülhet fel a kérdés, hogy ki viseli a felelősséget az egyes döntés átláthatóságát illetően. (Európai Tanács, dátum nélk.) A kiberbiztonsági fenyegetésekkel is érdemes lehet számításba vennünk, hiszen a mesterséges intelligencia sebezhetőségét kihasználhatják akár egy adatmérgezéssel is, ami máris torzíthatja a kapott eredményeinket. (Kocacevic, Radenkovic, & Nikolic, 2024)

5. Adatvédelem és adatbiztonság

A mesterséges intelligencia alkalmazása a kockázatelemzésben és az adatvédelem kérdésköre elválaszthatatlanul összefüggésben áll egymással. Az AI rendszerek hatékonyságához ugyanis rengeteg adatra van szükség, amelyek sokszor szenzitív személyes adatokat is magukba foglalnak. Az ilyen típusú adatgyűjtés, adattárolás és adatkezelés számos komoly adatvédelmet érintő kockázatot vet fel, amelyeket a GDPR rendelet szemléletében igen komolyan kell vennünk. Az ezen adatokat feldolgozó szoftverek esetében kulcsfontosságú, hogy megfelelő adatkezelési gyakorlatokat, valamint védelmi intézkedéseket alkalmazzunk a bizalom megőrzése érdekében. Mindez magába foglalja a hozzáférések szabályozását, rendszeres biztonsági teszteléseket és a személyes adatok anonimizálását is. Ebben a fejezetben alaposan szeretném bemutatni a konkrét követelményeket mind az adatkezelés mind az adatbiztonság tekintetében.

5.1. Személyes adatok védelme

A GDPR 5. cikke alapján az adatkezelés célhoz kötötten, jogszerűen és átláthatóan gyűjthet és dolgozhat fel adatot. Ez az elv az adatminimalizálás betartását igyekszik elérni, mivel a mesterséges intelligencia modellek általában nagy mennyiségű adatokkal dolgoznak. Ez az intézkedés biztosítja az érintettek jogainak védelmét és igyekszik csökkenteni az adatvédelmi

kockázatok jelenlétét. (Bereczki, 2024) Annak érdekében, hogy biztosítani tudjuk a személyes adatok védelmét szükséges technikai és szervezési intézkedéseket alkalmazniuk az AI rendszerek fejlesztőinek és üzemeltetőinek. Erre jó megoldás lehet például az adatok titkosítása, a hozzáférések szigorú szabályozása is. Ezekkel az intézkedésekkel biztosíthatjuk az adatok bizalmasságának, rendelkezésre állásának lehetőségét. (Microsoft, 2025) A már említett GDPR 22. cikke értelmében a mesterséges intelligencián alapuló rendszereknek is lehetőséget kell biztosítaniuk az emberi beavatkozásra, hogy egyes döntéseket felül tudjanak vizsgálni. Erre azért van szükség, hogy az ügyfelek ne kizárólag automatizált adatkezelésen alapuló döntéshozatallal álljanak szemben. (Bereczki, 2024) Az adatvédelmi hatásvizsgálat (data protection impact assessment, DPIA) a GDPR 35. cikke szerint akkor alkalmazandó, amikor az adatkezelés „valószínűsíthetően magas kockázattal jár a természetes személyek jogaira és szabadságaira nézve”. Ennek a vizsgálatnak a célja az adatkezelés jellegi feltárása, szükségességének vizsgálata. (Nemzeti Adatvédelmi és Információszabadság Hatóság, dátum nélk.) Az alábbi esetekben különösen fontos a vizsgálat elvégzése: természetes személyekre vonatkozó egyes személyes jellemzők automatizált gyűjtése és értékelése, nyilvános helyek nagymértékű, módszeres megfigyelése. A NAIH (Nemzeti Adatvédelmi és Információszabadság Hatóság) elvárása szerint az alábbi esetekben kell adatvédelmi hatásvizsgálatot folytatni: „egy természetes személy biometrikus adatainak kezelése módszeres megfigyelésre irányul, ha kiszolgáltatott helyzetben lévő érintettekkel – különös tekintettel a gyermekekre, munkavállalókra, idős, mentális betegségben szenvedőkre – kapcsolatos biometrikus adat kezelése történik, az adatkezelés egy természetes személy genetikai adatainak egyéb különleges adatokhoz vagy fokozottan személyes jellegű adatokhoz történő hozzákapcsolásával jár, egy természetes személy genetikai adatai kezelésének célja a természetes személy értékelése vagy pontozása”. (Kulcsár, 2018) A vizsgálat elemeiként a tervezett adatkezelési műveletek részletes leírását, az adatkezelés céljainak ismertetését, az érintett jogait és szabadságait érintő kockázatok vizsgálatát, valamint a kockázatok kezelését célzó intézkedések bemutatását tartjuk számon. (Adatvédelmi rendelet, dátum nélk.) Az alapos átvizsgálásnak köszönhetően a mesterséges intelligencia alapú rendszerek könnyebben azonosítani tudják az adatkezelési kockázatokat, továbbá képesek megállapítani melyek a releváns, szükséges intézkedések azok kezelésére. (Kumar Murakonda & Shokri, 2020)

5.2. Társadalmi hatások

A mesterséges intelligencia mind a társadalom, mind a gazdaság felépítését formálja át. A korábbiakban már ismertettem a mesterséges intelligencián alapuló rendszerek előnyeit,

azonban az AI alkalmazása komoly munkaerőpiaci, etikai kérdéseket is felvethet. Ebben a fejezetben szeretném feltérképezni a mesterséges intelligencia hatásait különösen kiemelve a munkahelyek automatizálását, valamint a gazdasági növekedésre vonatkozó hatásokat.

A mesterséges intelligencia számos előnyének köszönhetően forradalmasíthatja a munkavégzést az egyes ágazatokban, így az olyan feladatok, amelyek eddig emberi beavatkozással, kézzel készültek ma már robotok által elvégezhetőnek tekinthetjük. (Tölgyes, 2024) Az automatizálás így alapjaiban alakíthatja át a jelenlegi világgazdaságot és munkaerőpiacot. A McKinsey & Company, egy globálisan elismert tanácsadó vállalat 2018-as tanulmányában már kimutatta, hogy a munkafeladatok közel 49 százaléka már technikailag automatizálható is. (Fine, és mtsai., 2018) A vállalatoknak ez természetesen jelentős költségcsökkentést jelenthet, azonban a munkavállalói oldalról tekintve aggasztó lehet a technológia fejlődése, ugyanis ez a technológiai előrelépés azt is jelentheti, hogy kevésbé van szükség az emberi munkaerőre. Ez a probléma leginkább az alacsonyabb képzettséget igénylő munkákat veszélyezteti, ahol a munkavégzés fő eleme közé tartoznak a rutinfeladatok. Természetesen az ilyen típusú munkát végző embereknél fennáll a lehetőség arra, hogy átképzéssel a mesterséges intelligencia által teremtetett új munkát vállaljanak azonban a vállalatok nem igazán figyelnek oda az alkalmazottjaik átképzésének támogatására. (Napi érdekes, dátum nélk.) A Built In (kezdő vállalkozások és technológiai cégek számára kínál lehetőséget állások meghirdetésére) egy 2024-es bejegyzésükben ezt a kérdéskört vizsgálta meg alaposan, hogy valójában mely szakmák esetében elképzelhető, és melyek esetében nem, hogy a mesterséges intelligencia segítségével automatizálhatóak az egyes feladatok, pozíciók. Az ügyfélszolgálati munkatársak esetében a már korábban is említett chatbotok, virtuális asszisztensek segítségével a vállalatok képesek egy éjjel-nappal elérhető ügyfélszolgálat biztosítására jelentős emberi beavatkozás igénye nélkül. A logisztikai ágazatokban is megfigyelhető az átalakulás, a raktárakban alkalmazott robotok segítségével, valamint az önvezető járművek alkalmazásával csökkenthető mind a raktári dolgozók, mind a hivatásos teherautósofőrök iránti kereslet. A technológia az informatikai pozíciókat is veszélyeztetheti, különös tekintettel a kezdő programozást igénylő munkakörökre. Ez azért is lehet, mert egyes mesterséges intelligencia modellek, például a ChatGPT vagy a GitHub Copilot már képesek alapvető kódolási feladatok elvégzésére. A kutatási, valamint az elemzési területeket is érintheti ez a veszély, az adatfeldolgozás gyorsaságából, valamint azok vizualizálásából eredően. A jogi asszisztensek által végzett háttérmunkák szintén modellezhetőek az algoritmusok segítségével. Ez a tendencia megfigyelhető a kreatív tartalomírás vagy a grafikai

tervezés területén is, szövegvázlatok létrehozásával, valamint automatikus képgenerálással. (Urwin, 2024) Az alábbi táblázatban szemléltetem az egyes munkakörökre való mesterséges intelligencia okozta hatásokat:

Munkakör	AI hatása (kulcsszavak)
Ügyfélszolgálati munkatárs	Chatbot, automatizálás, önkiszolgálás
Autó- és tehergépkocsi-vezető	Önvezetés, fuvarautomatizálás, sofőrcsökkenés
Programozó	Kódkiegészítés, generatív AI, alapszintű programozás
Kutatási elemző	Adatelemzés, mintafelismerés, vizualizáció
Jogi asszisztens	Dokumentumkezelés, jogi kutatás, riportgenerálás
Raktári/feldolgozóipari munkás	Robotika, gépi látás, logisztikaautomatizálás
Pénzügyi kereskedő	Algoritmus, gyors elemzés, trendjósítás
Utazási tanácsadó	Ajánlórendszer, keresőmotor, virtuális túrák
Tartalomíró	Szöveggenerálás, ötletelés, sablonos tartalom
Grafikus	Képgenerálás, AI-design, vizuális automatizáció

ábra 3 Mesterséges intelligencia hatása az egyes munkakörökben

Összességében tehát igaz az állítás, hogy a mesterséges intelligencia nem kizárólag az alacsony képzettséget igénylő, vagy fizikai munkát képes kiváltani. Ez rengeteg új kérdést vet fel a munkaerőpiacon, többek között a munkahelyek jövőjéről, az emberi és gépi munka arányáról, valamint az átképzés szükségességéről. A mesterséges intelligencia fokozza az emberek aggodalmát a munkahelyek elvesztésével kapcsolatban, az AI technológiák alkalmazása esetén a munkáltatónak szükséges figyelembe vennie a társadalmi hatásokat is, ezzel ugyanis csökkenthető az intézkedések okozta frusztráció az alkalmazottak szemében. (Alphasonic, 2024)

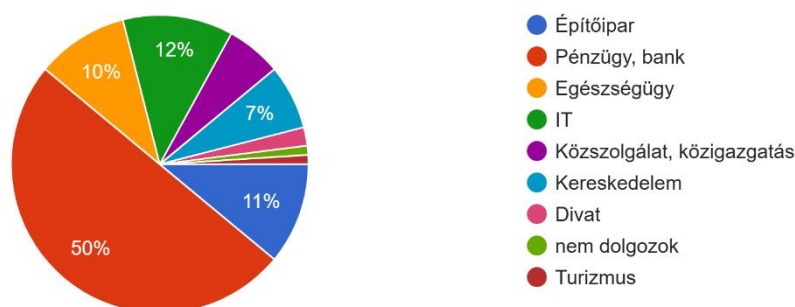
6. A kérdőív elemzése

A kérdőív célja az volt, hogy feltérképezzem a válaszadók hogyan viszonyulnak a mesterséges intelligencia alkalmazásához a pénzügyi kockázatelemzésben. Ennek érdekében egy strukturált kérdőíves felmérést készítettem, amelyet a témában érintett vagy érdeklődő személyek töltöttek ki elsősorban. Az alábbiakban a válaszok alapján leszűrt következtetéseket, összefüggéseket szeretném bemutatni.

6.1. Demográfiai adatok

A kérdőív első szakasza a demográfiai adatokat vizsgálta. A tesztben demográfiai adatok szerint 61%-ban nők, 39%-ban férfiak vettek részt. A legtöbb válaszadó a 18-25 éves korosztályból került ki 31%-kal, őket a 26-35 éves korosztály követi 26%-kal. Ebből következhet az, hogy a válaszadók legmagasabb iskolai végzettsége az egyetemi alapképzés (48%). A válaszadók 50%-a azonban a pénzügyi vagy a bankszektorban dolgozik, ez adódhat a kérdőív jellegéből is, tekintve, hogy talán számukra a legrelevánsabb a téma.

Ön melyik iparágban dolgozik?



ábra 4 Kitöltők aránya iparáganként

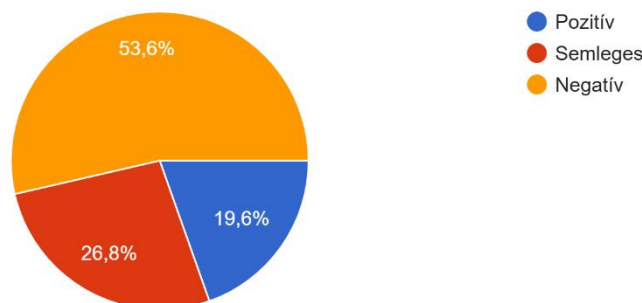
6.2. Kockázatok

A kérdőív második szakaszában a kockázatokat helyeztem a középpontba. Elsőként szerettem volna tudomást szerezni arról, hogy maga a kockázat szó milyen reakciót vált ki a személyekből. A válaszadók 53,6%-a negatív, 26,8%-a semleges és mindössze 19,6%-a

pozitív benyomásról tanúskodott.

Önnek milyen benyomása van a kockázat szó hallatán?

97 válasz

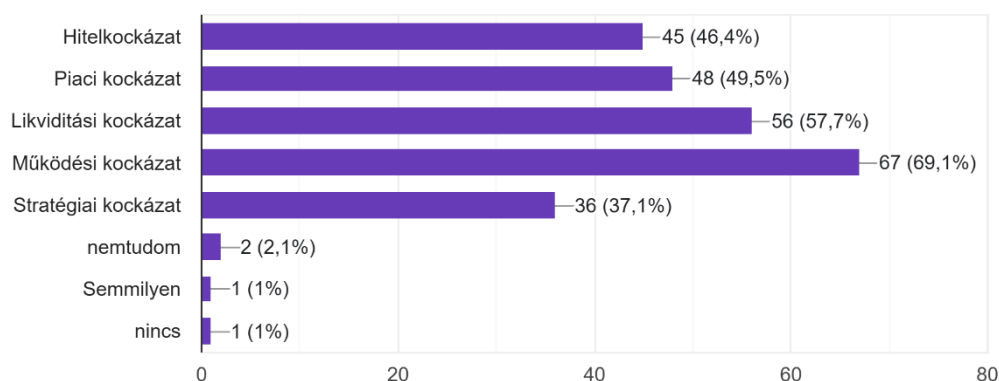


ábra 5 Vélemények megoszlása a kockázat szóról

Ami a kockázatok típusát illeti a legtöbben a működési kockázatot emelték ki, mint számukra releváns munkahelyi kockázat. Ennek az lehet az oka, hogy ez a napi munkavégzés során felmerülő hibákat, rendszerleállásokat is magába foglalja. Az ilyen kockázatok minden munkavállalót érintenek függetlenül attól, hogy pontosan milyen területen dolgoznak. A második leggyakrabban említett kockázat a likviditási kockázat volt, ez arra enged következtetni, hogy a munkavállalók tisztában vannak azzal, hogy a vállalat pénzügyi zavarai hatással lehetnek a működésre, a bérekre. Ez különösen érzékenyen érintheti az alkalmazottakat, hiszen hatással lehet biztonságérzetükre, hosszabb távon akár a vállalat fenntarthatósága is megkérdőjeleződhet. A harmadik leggyakrabban megjelölt kockázati típus a piaci kockázat, amely bizonytalanság a kamatkörnyezet ingadozásából, árfolyammozgásokból is eredhet. A prioritásként megjelölt kockázat jelezheti azt, hogy a pénzügyi szektor dolgozói kifejezetten érzékenyek az efféle külső tényezőkre. A válaszadók jelentős része megjelölte a hitelkockázatot is, amely arra utalhat, hogy a leginkább banki és pénzügyi szektorban dolgozók számára értelmezhető, nemfizetésből eredő veszélyek miatt jelölték meg ezt a típust. A válaszok alapján feltételezhető, hogy a válaszadók egy része aktívan foglalkozik hitelezéssel, vagy átlátja, hogy a megnövekedett ügyfélkockázatok a cég teljesítményében is megmutatkoznak. A stratégiai kockázat a kérdőív szerint alacsonyabb, de még mindig jelentőséggel bíró kockázattípus. Ez utalhat arra, hogy a munkavállalók úgy érzik jelenlegi munkahelyük nem a megfelelő irányba halad, vagy nem tudja kellőképpen felvenni a versenytársakkal a harcot. Ezt a kockázattípust valójában inkább a felsővezetés érzi, ebből adódhat az alacsonyabb megjelölés, ugyanakkor a stratégiai kockázat igencsak nagy szerepet játszik a vállalat hosszú távú működésében.

Ön szerint a kockázatelemzés mely típusai a leggyakoribbak az Ön munkájában? (több válasz is lehetséges)

97 válasz

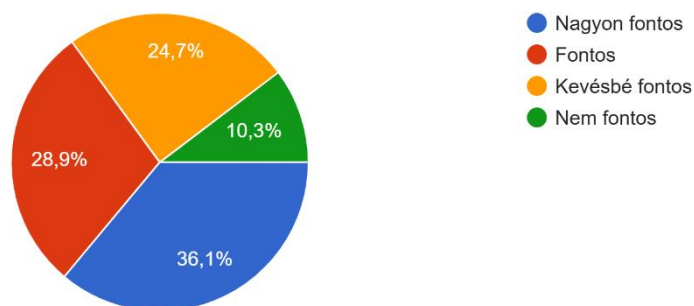


ábra 6 Kockázatok típusainak megoszlása

A válaszadók többsége 36,1%-a nagyon fontosnak találja a kockázatkezelést a mindennapi munkavégzés során, míg 28,9% fontosnak ítélte meg. Ez azt jelenti, hogy a válaszadók 65%-ának mindennapjaiban jelentős szerepet játszik a kockázatkezelés. Az eredmények azt igazolják, hogy a munkavállalók többsége tudatosan viszonyul a kockázatokhoz.

Mennyire tartja fontosnak a kockázatkezelést a mindennapi munkájában?

97 válasz



ábra 7 Kockázatkezelés fontossága

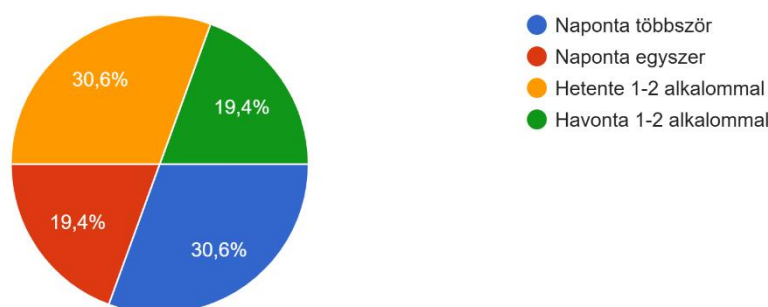
6.3. Mesterséges intelligencia

A kérdőív harmadik szakasza a mesterséges intelligenciáról szólt. Elsőként azt vizsgáltam, hogy használnak-e egyáltalán a munkavállalók mesterséges intelligenciát a munkájuk során. A válaszadók 64,9%-a nem használja, azonban 35,1%-uk alkalmazza ezt a technológiát. Az alkalmazás gyakoriságát illetően megoszlottak a vélemények. Egyenlő arányban, 30,6%-ban

vannak, akik naponta akár többször is alkalmazzák, azonban ehhez az arányhoz tartoznak azok is, akik hetente 1-2 alkalommal hívják segítségül. Meglepően azok aránya is megegyezik, akik naponta egyszer, vagy havonta 1-2 alkalommal alkalmaznak mesterséges intelligencia alapú eszközöket a munkájukban. Ez a statisztika adódhat abból, hogy a kitöltők kevésbe tartják magukat magabiztosnak a mesterséges intelligencia alkalmazásában. A kitöltők mindössze 15,5%-a érzi magát magabiztosnak a mesterséges intelligencia alkalmazása közben. Ez az alacsony arány arra utalhat, hogy a technológiai újdonságok még nem épültek be az emberek mindennapjaiba, még nem részei a rutinnak. A mesterséges intelligencia használatával kapcsolatos bizonytalanság fokozhatja az érzékelt működési és stratégiai kockázatokat is, mivel a technológia nemcsak lehetőséget, hanem komplexitást és potenciális hibaforrást is jelent a válaszadók szemében.

Amennyiben IGEN, milyen gyakorisággal?

36 válasz

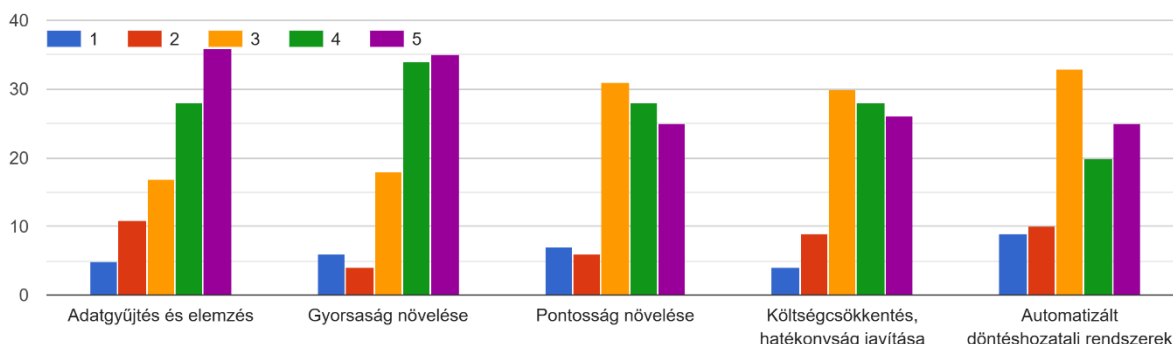


ábra 8 Mesterséges intelligencia alkalmazásának gyakorisága

6.4. AI-alapú megoldások hatása a kockázatelemzésre

A negyedik szakaszban a mesterséges intelligencia alkalmazására tértem ki a kockázatelemzésben. A kitöltők kicsivel több mint a fele még nem hallott a mesterséges intelligencia alapú megoldásokról a kockázatelemzésben, ennek ellenére készségesen megosztották véleményüket a témáról. Egyértelműen kimutatható, hogy a legtöbben az adatgyűjtésben és elemzésben, a gyorsaságban látnak fantáziát, ilyen típusú feladatokban hívnák leginkább segítségül a mesterséges intelligenciát. Szkeptikusabban álltak azonban ahhoz, hogy a pontosságot, a hatékonyságot és az automatizált döntéshozatalt is elősegíthetik a mesterséges intelligencia alapú rendszerek. Ez azt jelezheti, hogy az AI-t leginkább támogató eszközként tudják elképzelni, illetve a korábban már vizsgált viszolygás is okot adhat ezekre az adatokra.

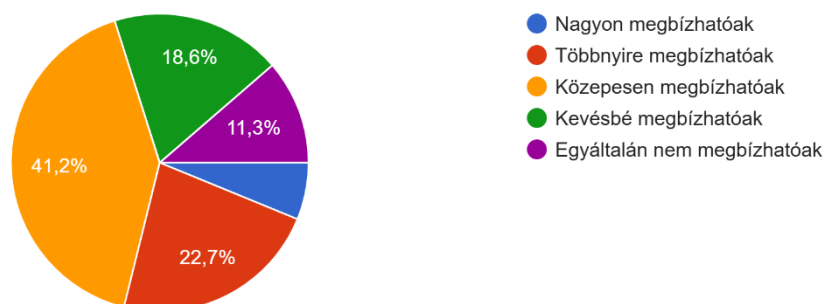
Mit gondol, milyen mértékben segíti az AI a következő szempontok mentén a kockázatelemzést? (1-től 5-ig skálán, ahol 1 - egyáltalán nem segít, 5 - nagyon sokat segít)



ábra 9 AI a pénzügyi kockázatelemzésben

A válaszadók többsége mérsékelt bizalommal kezeli a jelenlegi mesterséges intelligencia alapú rendszerek kockázatkezelési stratégiáit, a legtöbben közepesen megbízhatónak ítélték meg ezeket. A nagyon megbízható minősítést mindössze 6% adta meg, míg 22,7% szerint többnyire megbízhatóak. A megoszlás arra utal, hogy az AI-alapú kockázatelemzés még nem nyerte el teljes mértékben a felhasználók bizalmát – főként valószínűleg az átláthatóság, magyarázhatóság, vagy a technológiai komplexitás miatt.

Véleménye szerint a jelenlegi AI-alapú kockázatelemzési rendszerek mennyire megbízhatóak?
97 válasz



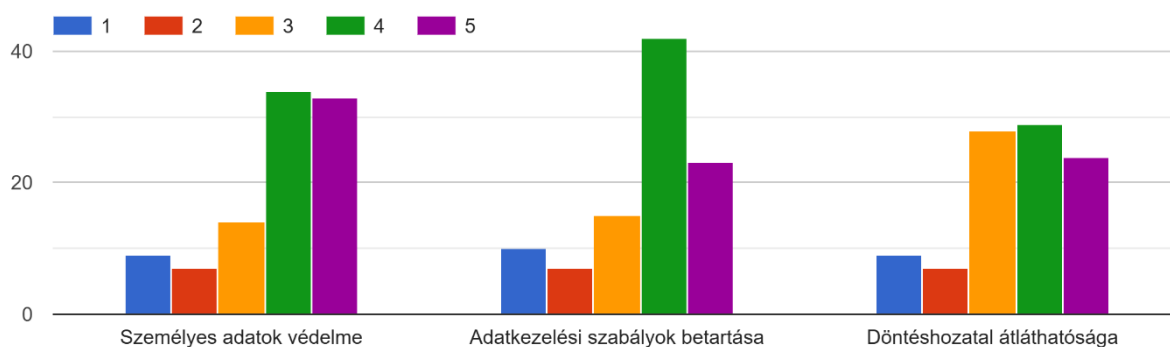
ábra 10 AI megbízhatósága

6.5. Etikai kihívások és adatvédelem

A kérdőív ötödik szakaszában az etikai kihívások és az adatvédelem kerültek a főszerepbe. Leginkább azt vizsgáltam, hogy mennyire tartják problémásnak a jelenlegi feltételek mellett a

mesterséges intelligencia etikus alkalmazását. A válaszadók kiemelten problémás területnek gondolják a személyes adatok védelmét, illetve az adatkezelési szabályok betartását. Ezzel szemben a döntéshozatal átláthatóságát kevésbé tartják veszélyesnek. Ez arra utal, hogy az adatbiztonság kiemelten fontos a kitöltők szemében. A döntéshozatal átláthatóságára kapott eredmény utalhat arra, hogy a kitöltők személyesen még nem találkoztak a fekete doboz jelenséggel, esetleg nem is kérdőjelezi meg, hogy a rendszerek biztosan pontos, naprakész adatokból dolgoznak-e.

Milyen mértékben tartja problémásnak az alábbi etikai kérdéseket az AI alkalmazásában? (1-től 5-ig skálán, ahol 1 - egyáltalán nem problémás, 5 - nagyon problémás)

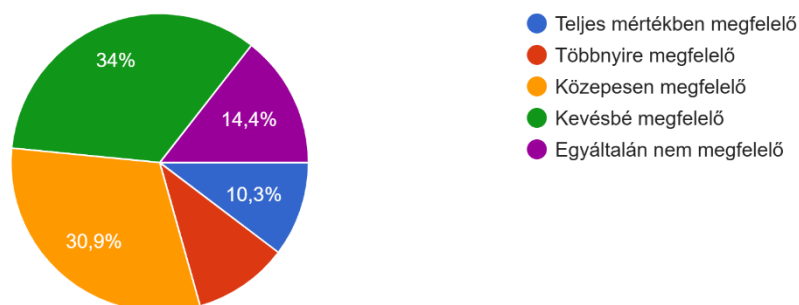


ábra 11 Etikai kérdések és az AI összeférhetősége

A válaszadók a jelenlegi mesterséges intelligencia szabályozási környezetét nem találja megfelelőnek annak etikus alkalmazásához. Mindössze 20,6% gondolja úgy, hogy a szabályozások teljes mértékben vagy többnyire megfelelőek, míg 30,9% közepesen megfelelőnek értékeli azokat. Ezzel szemben 34% kevésbé megfelelőnek, 14,4% pedig egyáltalán nem megfelelőnek ítéli meg a jelenlegi keretrendszert. Ez az eloszlás arra utal, hogy széles körű bizonytalanság és bizalmatlanság tapasztalható a szabályozások gyakorlati alkalmazhatóságát illetően.

Mennyire érzi úgy, hogy megfelelő szabályozások állnak rendelkezésre az AI-alapú kockázatelemzés etikus alkalmazásához?

97 válasz



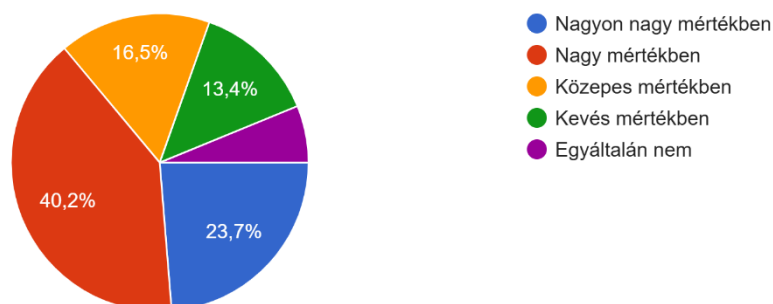
ábra 12 Vélemények megoszlása a szabályozásokról

6.6. Jövőbeli kilátások

A további kérdések már a jövőre irányultak. Arra voltam kíváncsi, hogy mit gondolnak a válaszadók, milyen mértékben alakíthatja át a pénzügyi szektort a mesterséges intelligencia térhódítása. A többség szerint a következő 5-10 évben alaposan átalakíthatja majd a pénzügyi kockázatelemzést az új technológia, összesen több, mint 63%- a a válaszadóknak optimistán állt a kérdéshez. Az eredmények azt tükrözik, hogy a szakmabeliek nagy változásokra számítanak a közeljövőben.

Ön szerint az AI technológia milyen mértékben fogja átalakítani a pénzügyi kockázatelemzést a következő 5-10 évben?

97 válasz



ábra 13 AI a következő 5-10 évben

A kérdőív nyílt kérdésében lehetőségük nyílt a válaszadóknak megosztani véleményüket, hogy pontosan milyen változásokra számítanak a mesterséges intelligencia alkalmazását

követően. A válaszokból több összefüggést is kiolvashatunk. A legtöbben az automatizációt, a munkafolyamatok gyorsulását emelték ki. Megfigyelhető azonban a félelem is az AI végett, ugyanis többen is megemlítették a munkahelyek megszűnését. Többen hangsúlyozták azt is, hogy csak és kizárólag a megfelelő szabályozások mellett alkalmazható biztonságosan a mesterséges intelligencia a mindennapi munkavégzés során. A válaszadók szerint a mesterséges intelligencia azonban nem képes helyettesíteni az emberi munkaerőt, csupán elősegíti, támogatja a munkavégzést.

Ön szerint milyen trendek várhatók a pénzügyi szektorban az AI fejlődésének köszönhetően?

18 válasz

automatizált rendszerek
A megfelelő szabályrendszer kialakítása esetén az adatgyűjtés és feldolgozás nagymértékű automatizálása, a döntéshozatal bizonyos szintű automatizálása valósulhat meg.
Költségcsökkentés, nagyobb hatékonyság, gyorsabb döntéshozatal
Felgyorsult adatrendszerzés
Úgy vélem az AI nem teljesen fogja a humán erőforrást helyettesíteni a pénzügyi szektorban, azonban bizonyos területek folyamatait jelentősen meg fogják tudni gyorsítani.
Folyamatok gyorsulnak fel, létszámleépítés.
call centerek megszűnése, automatizált folyamatokat kiváltja
Az a GDPR betarthatóságától függ.

ábra 14 Nyílt kérdés a trendekről

6.7. Összegzés

A válaszokból kiderül, hogy a mesterséges intelligencia jelentős hatással lehet a pénzügyi szektorra a jövőben. A leginkább várt funkciói közül az automatizáció, a gyorsaság kiemelkedő helyen szerepelnek, azonban az adatvédelem és az etikai szabályozások szigorú feltételei mellett. Az AI szerepét alapvetően tehát szkeptikusan értékelték a kitöltők, ugyanakkor nyitottak a technológia által nyújtott újításokra a pénzügyi kockázatelemzésben.

7. Esettanulmányok

A korábbiakban megismerhettük a mesterséges intelligencia nyújtotta lehetőségeket, illetve annak nehézségeit, árnyoldalait is. Az alábbi esettanulmányokkal szeretném szemléltetni a

mesterséges intelligencia használatának sikerességét, illetve a kudarcait és azok tanulságait is. A következő alfejezetekben részletesen tárgyaljuk tehát a sikeres, sikertelen projekteket, valamint elemezzük is az eseteket.

7.1. Sikeres mesterséges intelligencia alapú kockázatelemzési projektek

7.1.1. Shawmut Design and Construction

A Shawmut Design and Construction, egy bostoni székhelyű építőipari vállalat, évi kétmilliárd dolláros árbevétellel rendelkezik, közel harmincezer alkalmazottat, alvállalkozót és beszállítót foglalkoztatnak napi szinten. A vállalat már 2017 óta alkalmazza a mesterséges intelligenciát a munkavédelmi kockázatok előrejelzésére, valamint a munkavállalók biztonságának növelésére érdekében. (Villano, 2025)

A társaság mesterséges intelligencia rendszere rengeteg adatforrást alkalmaz, például az Országos Meteorológiai Szolgálat időjárás előrejelzéseit, valamint a személyi változásokat is elemzi, ezzel megismerik a potenciális veszélyeket. Vegyünk egy példát: a mesterséges intelligencia rendszer feltételezhető kockázatnak jelöli meg azt az esetet, ha egy munkaterületen a vezetők száma nem arányosan növekszik a hirtelen megnőtt új dolgozók számával. Ez a vezetőségnek azért fontos, mert így időben közbeavatkozhatnak, esetlegesen újra gondolhatják a feladatokat a lehető legpontosabb, leghatékonyabb munkafolyamatok kialakítása érdekében. (Villano, 2025)

A vállalat a nagy világjárvány, a COVID-19 idejében GPS-alapú mesterséges intelligencia rendszert vezetett be, ezzel biztosítva az előírt távolságtartást. A minden munkavállaló mobiltelefonján megtalálható GPS nyomkövető rendszer segítségével a mesterséges intelligencia automatikusan riasztott, amikor az alkalmazottak 6 lábnaál (körülbelül 2 méter) közelebb kerültek egymáshoz. Ezt a technológiát a járvány után is használják a magasban az állványok helyes használatának ellenőrzésére is. (Villano, 2025)

A munkavállalók nyomon követésével azonban felmerültek etikai kérdések is. A vállalat azt a megoldást választotta, hogy minden munkavállalói adatot anonimá tette, így megőrizték az alkalmazottak magánéletét és bizalmát egyaránt. Emellett a vállalat által kezelt adatgyűjtési folyamatok átláthatóak, így biztosítják, hogy a nyomon követésen ténylegesen csak és kizárólag a biztonság megőrzése érdekében történjen. (Villano, 2025)

Ami a jövőbeli kilátásokat illeti, a vállalat tervezi bővíteni a mesterséges intelligencia programjait, kifejezett figyelmet fordítva a valós idejű válaszadási képességek fejlesztésére.

Céljuk egy olyan rendszer megalkotása, amely valós időben képes figyelmeztetni a vezetőket, ha egy alkalmazottuk veszélyes területre lép, ez persze további kihívásokkal jár az adatok átláthatósága és megbízhatósága okán. (Villano, 2025)

Összefoglalva a Shawmut Design and Construction jól bemutatja, hogy a mesterséges intelligencia hatékonyan alkalmazható a munkavédelmi területeken, miközben az adatvédelmet és az etikai szempontokat is figyelembe vesszük.

7.1.2. RAZE Finance

A RAZE Finance egy fintech vállalat, amely ötvözi a hagyományos pénzügyi szolgáltatásokat az innovatív technológiával. Elsősorban a valós eszközök tokenizálására specializálódtak, ezzel teszik lehetővé a vállalatok számára, hogy digitálisan kezelhessék és értékesíthessék eszközeiket. Ezzel a lehetőséggel javíthatják a tőkeemelési folyamatokat, növelhetik a likviditást. (RAZE, 2025) A cég nem hagyományos bankként működik, ugyanis decentralizált pénzügyi megoldásokat és blokklánc technológiát kínál a pénzügyi szolgáltatások esetében, ezzel nagy mértékben hozzájárulva a pénzügyi szektor digitalizációjához, a hagyományos rendszerek átalakításához. (Raze , 2025)

AZ RTS Labs kifejezetten egy mesterséges intelligencia tanácsadó vállalat, akik segítenek feltárni a mesterséges intelligencia adta lehetőségeket. A bank az említett cég szakértőivel együttműködve kifejlesztett egy olyan mesterséges intelligencián alapuló rendszert, amely kifejezetten a működési hatékonyságot, a szabályozási megfeleléseket és a csalások megelőzését helyezte előtérbe. A RAZE banking kockázatkezelési módszerei hagyományos működésük végett nem tudták kezelni a gyorsan változó környezetet, valamint a csalási kísérletek számának növekedését. Az RTS Labs által elkészített mesterséges intelligencia rendszer alkalmazza mind a gépi tanulási modelleket, mind a prediktív analitikát, ezzel lehetővé téve, hogy valós időben tudják elemezni az ügyfélviselkedést, a szabályoknak való megfelelést, valamint a tranzakciókat. (RTS Labs, 2024) Az említett rendszer bevezetését követő három hónapban a bank 45%-os csökkenést realizált a csalást érintő tranzakciók számában. A szabályoknak való megfelelési mutató hatékonysága is 20%-kal javult, az operatív hatékonyság esetében ez 30%-kal nőtt. Ez a mértékű hatékonyságnövelés hírnéverősödést és jelentőségteljes csökkenést eredményezett a költségek terén. (RTS Labs, 2024) A Raze Banking kiemelkedő figyelmet fordított az etikai és adatvédelmi szempontokra is a mesterséges intelligencia alkalmazása során, ez ugyanis kulcsfontosságú volt a bank számára.

A pénzügyi intézmények tapasztalatai igazolják, hogy nem kizárólag a kockázatkezelésben lehet segítségünkre a mesterséges intelligencia, hozzájárulhat mind a szabályozások javításához, mind az operatív hatékonyság növeléséhez. Nagyon fontos azonban megemlítenünk, hogy a pénzügyi intézmények számára kiemelten fontos, hogy megfelelően kezeljék ügyfeleik adatait, hogy azok mind etikai, mind a szabályozási kritériumoknak megfeleljenek és biztosítani tudják az átláthatóságot és az elszámolhatóságot. (RTS Labs, 2024)

7.1.3. Upstart

Az Upstart egy vezető fintech innovatív vállalat, amely a hitelkockázat-értékelés esetében hívja segítségül a mesterséges intelligencia nyújtotta lehetőségeket, különösen a fiatalok esetében. A mesterséges intelligencia alkalmazásával az Upstart képes a gyorsabb és pontosabb hitelképesség értékelésére. (Redress compliance, 2025)

A mesterséges intelligencia bevezetését megelőzően a vállalat számos kihívással nézett szembe. Az egyik fő nehézség az volt, hogy korlátozottan jutottak hitelekhez a fiatal, hagyományos hiteltörténettel nem rendelkező ügyfelek. A másik fő kérdés a nemfizetési kockázat, ugyanis a tradicionális modellek nem képesek előre jelezni a FICO pontszámokon túli hitelkockázatot. A folyamataikat lassította az is, hogy gyakran manuális közbeavatkozást igényeltek a hitelbírálatok, illetve a hagyományos modellek alkalmatlanok voltak a nem hagyományos pénzügyi háttérrel rendelkezők hitelképességét. Ezen problémák kezelésére bevezettek egy mesterséges intelligencia alapú hitelkockázat értékelési programot, hogy szélesebb körben is elérhessék a potenciális ügyfeleiket. (Redress compliance, 2025)

A vállalat gépi tanulási technikák segítségével méri fel a hitelképességet. A modell több ezer adat elemzésével minél pontosabb és jobb döntéseket segít elő, továbbá csökkenti a hitelezők kockázatvállalását is. A vállalat AI rendszerei különböző pénzügyi veselkedési normákat elemeznek, mint például az ingatlanok bérleti díja, a közüzemi számlák, vagy a munkatapasztalat. Ezzel az elemzési módszerrel meggyőződhetünk a pénzügyi stabilitásról, amelyeket a hagyományos módszerek esetleg nem vennének számításba, ezzel javítható a hitelek jóváhagyási aránya is az alulbankolt személyek számára. (Redress compliance, 2025)

A hitelek jóváhagyásához prediktív elemzést végeznek, ezzel pontosabban előre tudják jelezni a hitelviszafizetés valószínűségét. A vállalat mesterséges intelligencia modellje folyton tanul a korábbi hitelteljesítményekből, így javítva a kockázatértékelési modellek teljesítményét, amely már a döntéshozatalnál óriási szerepet játszik. A prediktív modellezéssel biztosítják továbbá a nemteljesítési arányok csökkentését, illetve a hitelezők veszteségeinek

minimalizálását, hiszen a magas kockázatú kérelmezőket alaposabban képesek megszűrni. (Redress compliance, 2025) A mesterséges intelligencia képes a kérelmezőhöz igazítani mind a kamatlábakat, a jóváhagyási feltételeket és magukat a hitelfeltételeket is. Ezáltal a hitelfelvevők lehetőségeik szerint a legjobb ajánlatokhoz juthatnak, miközben a hitelezők csökkenthetik a kockázatokat és egy skálázhatóbb, hatékonyabb hitelezési folyamatot alakítanak ki. (Redress compliance, 2025)

A mesterséges intelligencia alkalmazásával az Upstart vállalat jelentősen javított a hitelezési folyamatai hatékonyságán és a hitelfelvevőkhöz való hozzáférésén. A vállalat a technológiának köszönhetően képes gyorsabban, pontosabban jóváhagyni a hitelek a nemteljesítési kockázatok minimalizálásával együttesen. (Redress compliance, 2025) Az alábbi táblázatban a mesterséges intelligencia bevezetésének hatásait mutatom be az egyes folyamatokra vonatkozóan.

Mutató	AI alkalmazása előtt	AI alkalmazása után
Hitel jóváhagyási arány	Alacsonyabb a FICO szigorú követelményei miatt	27%-os növekedés a hitelengedélyek számában
Alapértelmezett árfolyam	Magasabb a pontatlan kockázati előrejelzések miatt	16%-os csökkenés a nemteljesítéseknel
Hitelhozáférés az alulbankolt személyek esetében	Korlátozott számban	Kibővített hozzáférés alternatív adatokon keresztül
Kölcsönfeldolgozási idő	Kézi ellenőrzések lassították a jóváhagyásokat	Gyorsabb jóváhagyások automatikus kockázatelemzéssel
Működési költségek	Magas, az emberi beavatkozás végett	Alacsonyabb költségek mesterséges intelligencia által vezérelt automatizálással
Elfogultság a kreditpontozásban	Hagyományos modellekben jelen van	Csökkentett torzítás az AI-vezérelt értékelések révén

ábra 15 A mesterséges intelligencia bevezetésének hatásai az Upstart vállalat esetében

Az Upstart tervei között szerepel, hogy további fejlesztéseket végezzenek AI rendszereiken, melyek lehetővé teszik a még pontosabb és személyre szabottabb hitelbírálat, illetve a proaktív kockázatkezelést.

7.2. Kudarok és tanulságok

7.2.1. Diszkriminatív AI rendszerek és azok következményei

A mesterséges intelligencia alkalmazása a társadalmilag kifejezetten érzékeny területein rengeteg diszkriminatív eredményhez vezetett, az algoritmusok ugyanis a múlt előítéleteit, egyenlőtlenségeit is képesek újratermelni, ha azok az adathalmazok, amelyekből tanulnak szintén torzításokat tartalmaznak.

A PredPol, egy rendészeti előrejelző rendszer, amely a történelmi bűnözési adatok alapján határozza meg, hogy hol következhet be jelenleg is bűncselekmény. A rendszerről a későbbi elemzések után kiderült, hogy leginkább fekete és latin-amerikai emberek lakta városrészekbe irányította a rendőröket függetlenül attól, hogy az adott térségben valóban magasabb volt-e a bűncselekmények aránya. A Governing kutatása szerint a rendszer használata rontotta az érintett közösségek és a hatóságok közötti bizalmat. (Aaron, Surya, & Gilbertson, 2021)

Az amerikai börtönrendszerekben használt COMPAS nevű szoftvert is rengetegen kritizálták diszkriminatív eredményei végett. A szoftver ugyanis hasonlóképpen járt el a PredPol-hoz, mivel a fekete bőrű elítélteket gyakrabban sorolta magasabb kockázatúnak, annak ellenére, hogy a fehér bőrű raboknál jellemzőbb, hogy újra bűncselekményt kövessenek el a szabadulásuk után a ProPublica adatai szerint. (Angwin, Mattu, & Kirchner, 2016)

A diszkrimináció azonban nem kizárólag a büntetőjogban jelent meg, ugyanis az álláskeresés, a biztosítások ágazatában is rendszeresen előfordult, hogy a mesterséges intelligencia rendszerek hátrányosan megkülönböztették a nőket, az időseket és a kisebb etnikumú embereket. (Dastin, 2018) Az Amazon ebből az okból kifolyólag leállította a belső toborzórendszerét, ugyanis az algoritmus megtanulta, hogy a technológiai szerepekre főként férfiakat vettek fel korábban, ezért kevesebb pontszámot adott a női pályázóknak. (Dastin, 2018)

A fenti példák jól bemutatják, hogy a mesterséges intelligencia által kapott adatokat kritikusan kell szemlélnünk, de már az adatgyűjtés és adattisztítás fázisában sem árthat felülvizsgálni a rendszereket. Egyre inkább szorgalmazzák, hogy rendszeres és független auditálás alá essenek a rendszerek, így jobban szűrhetővé válnának az ilyen jellegű torzítások. (Groves, dátum nélk.)

7.2.2. AI projektek kudarcaránya

A mesterséges intelligencia projektek kudarcaránya rendkívül magas. A RAND Corporation kutatása szerint a projektek több, mint 80%-a kudarcba fullad, nem éri el a kitűzött célt. (Ryseff, F.De Bruhl, & J.Newberry, 2024)A probléma legtöbbször abból ered, hogy a vezetők nem egyeztetnek megfelelően a technikai szakemberekkel a célokról, így félreértések keletkezhetnek a működés során. (Ryseff, F.De Bruhl, & J.Newberry, 2024) Rendszeresen előfordul az is, hogy egy szervezet számára nem állnak rendelkezésre a megfelelő mennyiségű és minőségű adatok a hatékony modell képezéshez. Gyakran előfordul az is, hogy a szervezetek a legújabb technológiákkal kívánnak dolgozni, ezáltal azonban kevesebb fókusz helyeznek a valódi üzleti problémák megoldására. Természetesen rendelkeznie kell egy vállalatnak a megfelelő infrastruktúrával is, ideértve az adatkezelési rendszereket is. Ez azért kiemelten fontos, mert az alapvető technikai háttér nélkül a projektek nem tudnak hatékonyan működni. (Ryseff, F.De Bruhl, & J.Newberry, 2024) magas kudarcarányok azonban nem kizárólag anyagi károkat okoznak, a mesterséges intelligenciába vetett bizalmat is károsítja. Az S&P Global Market Intelligence 2025-ös bejegyzése alapján az AI projektek elhagyási aránya 17%-ról 42%-ra emelkedett mindössze egy év alatt. (Wilkinson, 2025)

Ahhoz, hogy sikeres AI projektet tudjunk véghez vinni elengedhetetlen a folyamatos kommunikáció a társterületek között, valamint, hogy közös célokat tűzzenek ki a szakemberek. (Ryseff, F.De Bruhl, & J.Newberry, 2024) A sikeres projektek alapja a megbízható, minőségi és releváns adatok halmaza, enélkül már a kezdetekben is hatalmas torzításokkal találkozhatunk a végeredményekben. (Ryseff, F.De Bruhl, & J.Newberry, 2024)A megfelelő technikai háttérrel, valamint a reális idő, erőforrás és hosszú távú elköteleződés együttesével sikeres mesterséges intelligencia projektet tudhatunk magunkévá. (Ryseff, F.De Bruhl, & J.Newberry, 2024)

8. Összefoglalás

A dolgozatom célja az volt, hogy átfogó képet alkosson a mesterséges intelligencia térnyeréséről a kockázatelemzésben, ismertesse az ebből eredő etikai és társadalmi kérdéseket. A tanulmány során világossá vált, hogy a mesterséges intelligencia alkalmazása hatalmas előrelépést jelenthet az elemzési pontosság, a hatékonyság és a gyorsaság szempontjából, azonban új problémákat is generál, különösen az adatok átláthatóságát illetően.

A mesterséges intelligencia rendszerek alkalmazása a bankok, biztosítók körében már eléggé elterjedt. Ezek a rendszerek ugyanis képesek az óriási adatbázisok villámgyors feldolgozására, továbbá képesek előre jelezni egy ügyfél fizetéseképtelenségének valószínűségét, valamint a csalásgyanús tranzakciókat. A technológia fejlődése azonban egy társadalmi, etikai kérdés is, így a szoftverfejlesztők felelőssége is egyre fontosabbá válik.

A kutatás részeként három hipotézist állítottam fel, az alábbiakban ezen hipotézisek eredményeit értékelem.

Az első hipotézisem szerint a hagyományos statisztikai modelleknél pontosabb előrejelzést nyújthatnak a mesterséges intelligencia alapú kockázatelemző rendszerek. A vizsgált szakirodalmi források egyértelműen alátámasztják az állítást. A mesterséges intelligencia modellek a gépi tanulás és a mélytanulás segítségével képesek olyan dinamikus összefüggések felismerésére, amelyek a hagyományos modellek esetében nem emelkednek ki az adathalmazokból. A kérdőíves kutatás eredményei alapján is helytálló a hipotézis, ugyanis a válaszadók többsége úgy vélte a mesterséges intelligencia az adatgyűjtésben, a gyorsaságban és a pontosság növelésében lehet segítségünkre a leginkább. Ugyanakkor az is nyilvánvalóvá vált, hogy a mesterséges intelligencia rendszerek a legjobban akkor alkalmazhatóak, ha az adatok minősége megfelel az elvárásoknak, illetve az egyéb kritériumoknak, ugyanis mindezek hiányában kialakulhat a fekete doboz modell.

A második hipotézisem szerint a mesterséges intelligencia alkalmazása bizalmatlanságot, átláthatatlanságot eredményez az ügyfeleink szemében. A kutatás során több forrásban is találtam példát, amelyek megerősítik ezt az állítást, kiemelkedő azonban a PredPol esete, ahol a rendszer diszkriminatív módon működött, ami a hírnevük romlását eredményezte. A kérdőíves kutatás is megerősítette ezt a félelmet, ugyanis a válaszadók mindössze 6%-a gondolta megbízhatónak a mesterséges intelligencia alapú rendszereket. Ez a bizalmatlanság fontos etikai kérdésekre ad okot, tekintve, hogy a pénzügyi döntések közvetlen hatással lehetnek az emberek életére, életminőségére.

A harmadik hipotézisem a mesterséges intelligencia rendszerek etikai kihívásaira helyezte a hangsúlyt, mivel gyakran személyes és szenzitív adatokat szükséges feldolgozniuk a rendszereknek. A kérdőíves kutatásból az derült ki, hogy a válaszadók szerint közepesen, vagy kevésbé megfelelőek a szabályozások a mesterséges intelligenciát illetően. Ez az eredmény igazolja azt, hogy a GDPR betartása mellett szükséges lehet egyéb, mesterséges intelligencia specifikus szabályok alkalmazására is.

Összességében megállapítható, hogy a mesterséges intelligencia alkalmazása a kockázatelemzésben nem kizárólag technológiai, hanem társadalmi és erkölcsi kérdéseket is felvet. A három hipotézis mindegyike beigazolódott, de egyúttal világossá vált, hogy az AI rendszerek sikeres és etikus alkalmazásához nem elég csupán a technikai kiválóság. Szükség van megfelelő szabályozásra, az átláthatóság biztosítására és a fejlesztők, döntéshozók etikai érzékenyítésére. Csak így biztosítható, hogy a mesterséges intelligencia valóban az emberi jólétet szolgálja, ne pedig új típusú kockázatok forrásává váljon.

9. Irodalomjegyzék

- Aaron, S., Surya, M., & Gilbertson, A. (2021). *Governing*. Forrás: <https://www.governing.com/security/crime-prediction-software-perpetuates-biases>
- Adatvédelmi rendelet*. (dátum nélk.). Forrás: <https://www.adatvedelmirendelet.hu/adatvedelmi-hatasvizsgalat/>
- Ahmad, Z. (2024. május). *Tabsgi*. Forrás: <https://www.tabsgi.com/ai-is-revolutionizing-financial-risk-assessment-learn-how/>
- Alphasonic*. (2024). Forrás: <https://www.alphasonic.hu/blog/a-mesterseges-intelligencia-15-legnagyobb-kihivasa-2024-ben/>
- Angwin, J., Mattu, S., & Kirchner, L. (2016). *Propublica*. Forrás: <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>
- Awati, R., & Yasar, K. (2024. október). *TechTarget*. Forrás: <https://www.techtarget.com/whatis/definition/black-box-AI>
- Bereczki, T. (2024). *Linked In*. Forrás: <https://www.linkedin.com/pulse/mesters%C3%A9gesintelligencia-rendszerek-%C3%A9s-gdpr-tamas-bereczki-9vqnf/>
- Crowdy.AI*. (dátum nélk.). Forrás: <https://crowdy.ai/hu/ai-chatbot-in-banking/>
- Dastin, J. (2018). *Reuters*. Forrás: <https://www.reuters.com/article/us-amazon-com-jobs-automation-insight-idUSKCN1MK08G/>
- Dr. Szikora, A., & Nagy, B. (dátum nélk.). *Mesterséges intelligencia a pénzügyi szektorban*. Forrás: <https://www.mnb.hu/letoltes/baksa-szikora-andrea-nagy-benjamin-ai-a-penzugyi-szektorban-final.pdf>
- Dujmovic, J. (2024. április). *MarketWatch*. Forrás: <https://www.marketwatch.com/story/financial-scammers-have-a-new-weapon-to-steal-your-money-ai-744eb000>
- Dyntell*. (dátum nélk.). Forrás: <https://dyntell.com/megoldasaink/a-valodi-ai-hasznalatba-vetele/prediktiv-analitika/>
- Erdős, M., & Mérő, K. (2010). *Pénzügyi közvetítő intézmények*. Budapest: Akadémiai Kiadó.
- Európai Tanács*. (dátum nélk.). Forrás: <https://www.consilium.europa.eu/hu/policies/benefits-and-risks-of-ai/>
- Fine, D., Havas, A., Solveigh, H., Jánoskúti, L., Kadocsa, A., & Puskás, P. (2018). *MC Kinsey*. Forrás: <https://www.mckinsey.com/~media/McKinsey/Locations/Europe%20and%20Middle%20East/Hungary/Our%20Insights/Transforming%20our%20jobs%20automation%20in%20Hungary/Automation-report-on-Hungary-HU-May24.pdf>
- Groves, L. (dátum nélk.). Forrás: <https://www.adalovelaceinstitute.org/project/new-york-algorithmic-bias-auditing/>

- Horváth S., A., Koltai, J., & Nádasdy, B. (2019). *Projekt- és Beruházásfinanszírozás Magyarországon*. Alinea Kiadó.
- Hovsepian, M. (2023). *Forbes*. Forrás: <https://www.forbes.com/councils/forbesfinancecouncil/2023/12/26/the-benefits-and-risks-of-ai-in-financial-services/>
- IBM. (2021. szeptember 22). Forrás: <https://www.ibm.com/think/topics/machine-learning>
- inf.u-szeged. (dátum nélk.). Forrás: https://www.inf.u-szeged.hu/~rfarkas/ML20/deep_learning.html
- Jasińska, D., & Puczyk, A. (2025. március). *Neontri*. Forrás: <https://neontri.com/blog/best-banking-chatbots/>
- Ken, D., Pramode, C., & Weiss, D. (2025. április). *Reuters*. Forrás: <https://www.reuters.com/legal/legalindustry/ai-agents-greater-capabilities-enhanced-risks-2025-04-22/>
- Kocacevic, A., Radenkovic, S., & Nikolic, D. (2024). *Cornell University*. Forrás: <https://arxiv.org/abs/2412.04495>
- KPMG. (2021. szeptember). Forrás: <https://kpmg.com/ae/en/home/insights/2021/09/artificial-intelligence-in-risk-management.html>
- KPMG. (2021. október). Forrás: <https://blog.kpmg.hu/2021/10/mesterseges-intelligencia-mint-a-kockazatkzezes-lehetseges-eszkoze/>
- Kulcsár, Z. (2018). *Adatvédelmi rendelet*. Forrás: <https://www.adatvedelmirendelet.hu/adatvedelmi-rendelet/a-hatasvizsgalat-lefolytatasanak-kotelezo-esetei/>
- Kumar Murakonda, S., & Shokri, R. (2020). *Cornell University*. Forrás: <https://arxiv.org/abs/2007.09339>
- Leitner, G., Singh, J., van der Kraaij, A., & Zsámboki, B. (2024). *European Central Bank*. Forrás: https://www.ecb.europa.eu/press/financial-stability-publications/fsr/special/html/ecb.fsrart202405_02~58c3ce5246.en.html
- Microsoft. (dátum nélk.). Forrás: <https://azure.microsoft.com/hu-hu/resources/cloud-computing-dictionary/what-is-machine-learning-platform>
- Microsoft. (2025). Forrás: <https://learn.microsoft.com/hu-hu/azure/cloud-adoption-framework/scenarios/ai/secure>
- Moody's. (2024. április). Forrás: <https://www.moody's.com/web/en/us/kyc/resources/insights/towards-interactive-smart-screening-with-generative-ai-in-kyc-workflows.html>
- Napi érdekes. (dátum nélk.). Forrás: <https://napierdekes.hu/a-mesterseges-intelligencia-ai-hatasa-a-mindennapi-életunkre-elonyok-kihivasok-es-etikai-kerdesek/>
- Nelson, D. (2024. március 20). *Unite.AI*. Forrás: <https://www.unite.ai/hu/what-is-natural-language-processing/>
- Nemzeti Adatvédelmi és Információszabadság Hatóság. (dátum nélk.). Forrás: <https://www.naih.hu/hatasvizsgalat>

OTP Bank. (dátum nélk.). Forrás: <https://www.otpbank.hu/portal/hu/rolunk/penzmosas-elleni-kuzdelem>

PWC. (dátum nélk.). Forrás: <https://www.pwc.com/m1/en/services/assurance/risk-assurance/documents/explainable-ai.pdf>

Raze , F. (2025). *Einpresswire*. Forrás: <https://www.einpresswire.com/article/797027181/raze-partners-with-plume-to-leverage-plume-chain-for-real-world-asset-innovation>

RAZE. (2025). Forrás: <https://raze.finance/>

Redress compliance. (2025). Forrás: <https://redresscompliance.com/ai-case-study-ai-powered-credit-scoring-risk-assessment-at-upstart/>

Rekettye, G., Töröcsik , M., & Hetesi, E. (2015). *BEVEZETÉS A MARKETINGBE*. Budapest: Akadémiai Kiadó.

Risk Officer. (dátum nélk.). Forrás: https://www.risk-officer.com/Reputation_Risk.htm

RTS Labs. (2024). Forrás: <https://rtslabs.com/ai-risk-management-finance>

Ryseff, J., F.De Bruhl, B., & J.Newberry, S. (2024). *Rand*. Forrás: https://www.rand.org/pubs/research_reports/RRA2680-1.html

SAP. (dátum nélk.). Forrás: <https://www.sap.com/hungary/products/data-cloud/cloud-analytics/what-is-predictive-analytics.html>

SEON. (dátum nélk.). Forrás: <https://seon.io/resources/guides/alternative-credit-scoring/>

Stefanoski, D., & Meier, K. (2024). *EY*. Forrás: https://www.ey.com/en_ch/insights/forensic-integrity-services/risks-and-benefits-of-generative-ai-in-financial-sector

Tiku, N. (2025. április). *Washington Post*. Forrás: <https://www.washingtonpost.com/politics/2025/04/22/ai-tools-mostly-fumble-basic-financial-tasks-study-finds/>

Tölgyes, L. A. (2024). *ICT Global*. Forrás: <https://ictglobal.hu/iparagi-megoldasok/az-mi-hatas-a-munkajovojere/>

Urwin, M. (2024). *Built In*. Forrás: <https://builtin.com/artificial-intelligence/ai-replacing-jobs-creating-jobs>

Villano, M. (2025). *Business Insider*. Forrás: <https://www.businessinsider.com/ai-for-worker-site-safety-in-construction-2025-4>

Visure Solutions. (dátum nélk.). Forrás: <https://visuresolutions.com/hu/blog/ai-%C3%A9s-g%C3%A9pi-tanul%C3%A1s-a-kock%C3%A1zatkezel%C3%A9shez/>

Wilkinson, L. (2025). *CIO Brief*. Forrás: <https://www.ciodive.com/news/AI-project-fail-data-SPGlobal/742590/>

Wolford, B. (dátum nélk.). *GDPR.EU*. Forrás: <https://gdpr.eu/what-is-gdpr/>

Yallo. (dátum nélk.). Forrás: <https://yallo.co/insights/industries/banking/ai-in-risk-management-predictive-analytics-for-financial-stability/>

10. Ábrajegyzék

ábra 1 Prediktív elemzés	13
ábra 2 Példák a digitális bankok által bevezetett mesterséges intelligenciára	17
ábra 3 Mesterséges intelligencia hatása az egyes munkakörökben	22
ábra 4 Kitöltők aránya iparáganként	23
ábra 5 Vélemények megoszlása a kockázat szóról.....	24
ábra 6 Kockázatok típusainak megoszlása	25
ábra 7 Kockázatkezelés fontossága	25
ábra 8 Mesterséges intelligencia alkalmazásának gyakorisága.....	26
ábra 9 AI a pénzügyi kockázatelemzésben.....	27
ábra 10 AI megbízhatósága	27
ábra 11 Etikai kérdések és az AI összeférhetősége	28
ábra 12 Vélemények megoszlása a szabályozásokról	29
ábra 13 AI a következő 5-10 évben.....	29
ábra 14 Nyílt kérdés a trendekről	30
ábra 15 A mesterséges intelligencia bevezetésének hatásai az Upstart vállalat esetében	34

11. Mellékletek

11.1. Kérdőív kérdései

- **Demográfiai adatok**
- **Az Ön neme: ***
 - Nő
 - Férfi
- **Az Ön életkora: ***
 - 18-25
 - 26-35
 - 36-45
 - 46-55
 - 56-65
 - 65+
- **Mi a legmagasabb iskolai végzettsége? ***
 - Általános iskola (8 osztály)
 - Érettségi bizonyítvány
 - Egyetemi alapképzés (BA/BSC)
 - Egyetemi mesterképzés (MA/MSc)
 - Egyéb:
- **Ön melyik iparágban dolgozik? ***
 - Építőipar
 - Pénzügy, bank
 - Egészségügy
 - IT
 - Közszolgálat, közigazgatás
 - Egyéb:

- **Kockázatok**
- **Önnek milyen benyomása van a kockázat szó hallatán? ***
 - Pozitív
 - Semleges
 - Negatív
- **Ön szerint a kockázatelemzés mely típusai a leggyakoribbak az Ön munkájában? (több válasz is lehetséges) ***
 - Hitelkockázat
 - Piaci kockázat
 - Likviditási kockázat
 - Működési kockázat
 - Stratégiai kockázat
 - Egyéb:
- **Mennyire tartja fontosnak a kockázatkezelést a mindennapi munkájában? ***
 - Nagyon fontos
 - Fontos
 - Kevésbé fontos
 - Nem fontos
- **Mesterséges Intelligencia**
- **Ön használ mesterséges intelligenciát a munkája során? ***
 - Igen
 - Nem
- **Amennyiben IGEN, milyen gyakorisággal?**
 - Naponta többször
 - Naponta egyszer
 - Hetente 1-2 alkalommal
 - Havonta 1-2 alkalommal
- **Mennyire érzi magabiztosnak magát az AI-technológiák használatában? ***
 - Nagyon magabiztos
 - Magabiztos
 - Kevésbé magabiztos
 - Nem magabiztos
- **AI-alapú megoldások hatása a kockázatelemzésre**
- **Hallott már AI-alapú megoldások alkalmazásáról a pénzügyi kockázatelemzésben? ***
 - Igen
 - Nem
- **Mit gondol, milyen mértékben segíti az AI a következő szempontok mentén a kockázatelemzést? (1-től 5-ig skálán, ahol 1 - egyáltalán nem segít, 5 - nagyon sokat segít) ***
 - Adatgyűjtés és elemzés

- Gyorsaság növelése
- Pontosság növelése
- Költségcsökkentés, hatékonyság javítása
- Automatizált döntéshozatali rendszerek
- Adatgyűjtés és elemzés
- Gyorsaság növelése
- Pontosság növelése
- Költségcsökkentés, hatékonyság javítása
- Automatizált döntéshozatali rendszerek
- **Véleménye szerint a jelenlegi AI-alapú kockázatelemzési rendszerek mennyire megbízhatóak? ***
 - Nagyon megbízhatóak
 - Többnyire megbízhatóak
 - Közepesen megbízhatóak
 - Kevésbé megbízhatóak
 - Egyáltalán nem megbízhatóak
- **Etikai kihívások és adatvédelem**
- **Milyen mértékben tartja problémásnak az alábbi etikai kérdéseket az AI alkalmazásában? (1-től 5-ig skálán, ahol 1 - egyáltalán nem problémás, 5 - nagyon problémás) ***
 - Személyes adatok védelme
 - Adatkezelési szabályok betartása
 - Döntéshozatal átláthatósága
 - Személyes adatok védelme
 - Adatkezelési szabályok betartása
 - Döntéshozatal átláthatósága
- **Mennyire érzi úgy, hogy megfelelő szabályozások állnak rendelkezésre az AI-alapú kockázatelemzés etikus alkalmazásához? ***
 - Teljes mértékben megfelelő
 - Többnyire megfelelő
 - Közepesen megfelelő
 - Kevésbé megfelelő
 - Egyáltalán nem megfelelő
- **Jövőbeli kilátások**
- **Ön szerint az AI technológia milyen mértékben fogja átalakítani a pénzügyi kockázatelemzést a következő 5-10 évben? ***
 - Nagyon nagy mértékben
 - Nagy mértékben
 - Közepes mértékben
 - Kevés mértékben
 - Egyáltalán nem
- **Ön szerint milyen trendek várhatók a pénzügyi szektorban az AI fejlődésének köszönhetően?**
 - Saját válasz

- **Van-e valamilyen további véleménye az AI alkalmazásáról a kockázatelemzésben?**
 - Saját válasz