



Mesterséges intelligencia módszerek alkalmazása az informatikai rendszerek biztonsági auditjában

Barta Gergő

Gödöllő

2021

A doktori iskola

megnevezése: Gazdaság- és Regionális Tudományok Doktori Iskola

tudományága: gazdaság és regionális tudományok

vezetője: Prof. Dr. H.c. Popp József
MTA levelező tag
Magyar Agrár- és Élettudományi Egyetem
Gazdaságtudományi Intézet - Intézetigazgató
Gazdaság- és Regionális Tudományi Doktori Iskola vezető

Témavezető: Dr. Pitlik László
Egyetemi Docens, Tanszékvezető
Kodolányi János Egyetem, Fenntartható Gazdaság Intézet
Informatika tanszék

.....
Az iskolavezető jóváhagyása

.....
A témavezető jóváhagyása

Tartalomjegyzék

1.	BEVEZETÉS	1
1.1.	A téma aktualitása, célkitűzései.....	1
1.2.	A dolgozat szerkezete	4
2.	SZAKIRODALMI ÁTTEKINTÉS	9
2.1.	Az interdiszciplináris kutatás tudományterületi kapcsolatrendszere.....	9
2.2.	Információrendszerek, biztonság és audit	10
2.2.1.	Az információ fogalmi megközelítései	10
2.2.2.	Információrendszerek megjelenése és biztonsági kérdései.....	11
2.2.3.	Információbiztonsági kontrollok auditálása.....	15
2.2.4.	Kontrollok tesztelése.....	17
2.2.5.	Az ISO:IEC 27001:2013 szabvány	19
2.2.6.	Az alfejezet összefoglalása	21
2.3.	Mesterséges intelligenciával ellátott modellfejlesztés.....	22
2.3.1.	A mesterséges intelligencia fogalmi megközelítései	22
2.3.2.	Gépi tanulás	27
2.3.3.	Gépi tanuló rendszerek típusai.....	29
2.3.4.	Gépi tanuló rendszerek fejlesztése	31
2.3.5.	A modellalkotás kihívásai és kockázatai	36
2.3.6.	Az alfejezet összefoglalása	39
2.4.	Gyanúgenerálás az információbiztonság kutatási területén.....	40
2.4.1.	A gyanúgenerálás fogalmi kerete.....	40
2.4.2.	Gyanúgenerálás a kibervédelemben.....	41
2.4.3.	Az alfejezet összefoglalása	45
2.5.	Hipotézisek felállítása a szakirodalmi áttekintés alapján	46
3.	ANYAG ÉS MÓDSZERTAN	47
3.1.	Adatgyűjtés	47
3.2.	Alkalmazott algoritmusok és statisztikai eljárások	50
3.2.1.	Döntési fa	51
3.2.2.	Neurális háló	52
3.2.3.	Adaptív Boosting	54
3.2.4.	Gradiens Boosting.....	56
3.2.5.	Kollaboratív szűrés	57

3.2.6.	Hasonlóságelemzés	57
3.2.7.	Pearson-féle korreláció	59
3.2.8.	Varianciaelemzés	60
3.3.	Gépi tanuló rendszerek kiértékelése	60
3.3.1.	Jóságmetrikák	60
3.3.2.	ROC-görbe és AUROC mutató	62
3.3.3.	PR-görbe és AUPRC mutató	63
3.3.4.	Keresztvalidáció	64
3.4.	Felhasznált eszközök és technológiai megoldások	64
3.5.	A kutatási célok és hipotézisek rendszere	65
4.	EREDMÉNYEK	67
4.1.	A kutatás során gyűjtött adatok leíró statisztikái.....	67
4.2.	Gyanúgenerálás információbiztonsági kontrollhiányosságok detektálására	75
4.2.1.	Adatfeldolgozás	76
4.2.2.	Felügyelt gépi tanuló modellek inicializálása.....	78
4.2.3.	Feltételezések modelljóságra a megalkotott felügyelt gépi tanuló modellek alapján....	79
4.2.4.	Felügyelt gépi tanuló modellek futtatása és kiértékelése	80
4.2.5.	Gyanúgenerálás felügyelet nélküli gépi tanuló módszerekkel	83
4.2.6.	Gyanúgenerálás hibrid megközelítésben	85
4.2.7.	Elemzések szűkített adathalmazokon	88
4.2.8.	Objektív modelljóság-beclés a generált modelleken anti-diszkriminatív eljárással	91
4.2.9.	Az alfejezet összefoglalása	94
4.3.	Genetikai potenciál keresése a gépi tanulás adatvagyonának redukált felhasználásával	95
4.3.1.	A javasolt innovatív keresési eljárás.....	95
4.3.2.	Genetikai potenciál keresése a rögzített adathalmazon	98
4.3.3.	Objektív modelljóság-beclés a javított modelleken anti-diszkriminatív eljárással	105
4.3.4.	Az alfejezet összefoglalása	106
4.4.	Modell-preferencia levezetése klasszikus tesztelési eljárások nélkül	107
4.4.1.	A javasolt eljárás.....	108
4.4.2.	A modellezés folyamata, alkalmazott gépi tanuló eljárások	110
4.4.3.	Feltételezések modelljóságra a megalkotott modellek alapján.....	112
4.4.4.	Objektumleíró tulajdonságok meghatározása	113
4.4.5.	Objektumleíró tulajdonságok irány-preferenciáinak ellenőrzése	114
4.4.6.	Modellek rangsorolása.....	116

4.4.7.	Eredmények értékelése a „minél kisebb, annál jobb” irány-preferencia nézetben	119
4.4.8.	Eredmények értékelése a „minél nagyobb, annál jobb” irány-preferencia nézetben..	122
4.4.9.	Objektív modelljóság-becslés a generált modelleken anti-diszkriminatív eljárással..	123
4.4.10.	Az alfejezet összefoglalása	124
4.5.	A dolgozat célkitűzéseinek teljesítése a SMART feltételrendszer alapján	125
5.	ÚJ ÉS ÚJSZERŰ TUDOMÁNYOS EREDMÉNYEK.....	127
6.	KÖVETKEZTETÉSEK ÉS JAVASLATOK.....	131
7.	ÖSSZEFOGLALÁS	135
8.	SUMMARY	137
9.	IRODALOMJEGYZÉK.....	139
10.	MELLÉKLETEK	153
	Köszönetnyilvánítás	186

Ábrák jegyzéke

1. ábra: A kockázatelemzés magas szintű folyamata.....	14
2. ábra: Az auditálás folyamata.....	16
3. ábra: Kockázatok, kontrollcélkitűzések és kontrollok kapcsolata	18
4. ábra: Audit technikák a bizonyosságszerzés mértékének rangsora szerint.....	19
5. ábra: A mesterséges intelligenciához szorosan kapcsolódó fogalmak népszerűsége	22
6. ábra: „Mesterséges Intelligencia” kulcsszót tartalmazó publikációk száma a Scopus-ban	22
7. ábra: A klasszikus programozás és gépi tanulás paradigmája	28
8. ábra: A gépi tanulás típusai.....	30
9. ábra: Gépi tanuló rendszerek fejlesztése a prediktív modellezés keretében	33
10. ábra: Osztályozási probléma különböző döntési határokkal.....	37
11. ábra: A döntési fa működési logikája.....	51
12. ábra: A neurális háló működési logikája.....	53
13. ábra: Az ABM működési logikája	55
14. ábra: ROC-görbe vizualizálása	62
15. ábra: A ROC-görbe kitüntetett pontjai.....	63
16. ábra: A PR-görbe vizualizálása.....	63
17. ábra: A keresztvalidáció működési logikája	64
18. ábra: A gyűjtött minta megoszlása iparáganként	67
19. ábra: A gyűjtött minta megoszlása audit típusonként	68
20. ábra: A 4.2. alfejezet rendszerezése	76
21. ábra: Az adatfeldolgozás folyamata.....	78
22. ábra: Felügyelt gépi tanuló algoritmusok performancia metrikái.....	81
23. ábra: Felügyelt gépi tanuló modellek ROC és PR-görbéi.....	81
24. ábra: Felügyelt gépi tanuló algoritmusok tanulási görbéi.....	82
25. ábra: A hibrid modellezés folyamatábrája	85
26. ábra: Az egyszerű és hibrid ABM algoritmus tanulási görbéi.....	87
27. ábra: Az egyszerű és hibrid GBM algoritmus tanulási görbéi.....	88
28. ábra: Az egyszerű és hibrid NN algoritmus tanulási görbéi	88
29. ábra: A javasolt, a genetikai potenciált kereső eljárás folyamatábrája	96
30. ábra: Modell-preferenciák levezetésére javasolt eljárás folyamatábrája	109
31. ábra: Felügyelet nélküli modellek irány-preferencia vizsgálata	116

Táblázatok jegyzéke

1. táblázat: A kutatás célkitűzéseinek értékelése a SMART feltételrendszer alapján	4
2. táblázat: Az ISO/IEC 27001:2013 szabvány kontrollterületei.....	20
3. táblázat: A mesterséges intelligencia definícióinak csoportosítása	23
4. táblázat: A dolgozatban ismertetett modellezési gyakorlatok szelektált forráshivatkozásai	50
5. táblázat: A korrelációs együttható iránya és erőssége	59
6. táblázat: Modellek értékelésére alkalmazott jóságmetrikák	60
7. táblázat: A kutatás során alkalmazott eszközök listája és leírása	64
8. táblázat: Célkitűzések, hipotézisek és alkalmazott módszerek rendszere	66
9. táblázat: Megállapítások, hatókörben lévő kontrollok száma, mutatói iparági megoszlásban	69
10. táblázat: Megállapítások, hatókörben lévő kontrollok száma, mutatói audit típusonkénti megoszlásban.....	70
11. táblázat: Megállapítások, hatókörben lévő kontrollok száma, aránya kontrollterületenkénti megoszlásban.....	71
12. táblázat: Top 10 kontrollkövetelmény az auditok hatókörében	72
13. táblázat: Top 10 legmagasabb megállapítással rendelkező kontrollkövetelmények	73
14. táblázat: Top 10 egy auditjelentésre eső megállapítások száma kontrollterületenként	74
15. táblázat: Felügyelt gépi tanuló eljárások technikai konfigurációi	79
16. táblázat: Felügyelt gépi tanuló algoritmusok performancia metrikái	80
17. táblázat: Felügyelet nélküli gépi tanuló algoritmusok performancia metrikái	84
18. táblázat: Egyszerű és hibrid ABM performancia metrikái	86
19. táblázat: Egyszerű és hibrid GBM performancia metrikái	86
20. táblázat: Egyszerű és hibrid NN performancia metrikái.....	86
21. táblázat: Egyszerű felügyelt modellek performancia metrikái audittípusonkénti megoszlásban	89
22. táblázat: Hibrid felügyelt modellek performancia metrikái audittípusonkénti megoszlásban....	89
23. táblázat: Felügyelt egyszerű módszerek performancia metrikái szűkített és a teljes adatvagyonon.....	90
24. táblázat: Felügyelt hibrid módszerek performancia metrikái szűkített és a teljes adatvagyonon	91
25. táblázat: Modelljóság-becslés anti-diszkriminatív eljárással	92
26. táblázat: Modell-jóság átlagok kategóriánként	93
27. táblázat: Az alappopuláció irány-preferenciáinak meghatározása.....	100
28. táblázat: Az alappopuláció súlyszerkezet az első iterációban	101
29. táblázat: Az újonnan beszűrt rekord értékei és az eddigi legjobb rekordtól történő elmozdulás minősítése	101
30. táblázat: A populáció súlyszerkezet a második iterációban	102
31. táblázat: A hibrid és a kereső eljárással javított ABM performancia metrikái	104
32. táblázat: A hibrid és a kereső eljárással javított GBM performancia metrikái	104
33. táblázat: A hibrid és a kereső eljárással javított NN performancia metrikái	104
34. táblázat: Modelljóság-becslés anti-diszkriminatív eljárással	105
35. táblázat: Felügyelet nélküli modellek leíró dimenzióinak összefoglaló táblázata.....	112
36. táblázat: Felügyelet nélküli modellek leíró statisztikáinak irány-preferencia értékelése hasonlóságelemzéssel	118
37. táblázat: Felügyelet nélküli modellértékelés összefoglaló táblázata az első szórás irány-preferencia nézetben	119

38. táblázat:	Felügyelet nélküli modellek értékeinek összefoglaló táblázata	120
39. táblázat:	Felügyelet nélküli modellek értékelő táblázata.....	120
40. táblázat:	Felügyelet nélküli modellek összefoglaló táblázata.....	121
41. táblázat:	Felügyelet nélküli modellértékelés összefoglaló táblázata a második szórás irány- preferencia nézetben.....	122
42. táblázat:	Felügyelet nélküli modellek értékeinek összefoglaló táblázata	122
43. táblázat:	Modelljóság becslés anti-diszkriminatív eljárással.....	124
44. táblázat:	A kutatás célkitűzéseinek teljesítése a SMART feltételrendszer alapján	125

Mellékletek jegyzéke

1. sz. melléklet:	A kutatás folyamatábrája.....	153
2. sz. melléklet:	Az ISO/IEC 27001:2013 A melléklet kontrollterületeinek rövid bemutatása	154
3. sz. melléklet:	A feldolgozott adatvagyon megoszlásai a teljes, tanuló-, valamint teszhalmazon audittípusonként és iparáganként.....	156
4. sz. melléklet:	Az egyszerű és hibrid modellek ROC és PR-görbéi	157
5. sz. melléklet:	Egyszerű és hibrid modellek jószágmetrikái grafikusan oszlopdiagrammon	158
6. sz. melléklet:	Egyszerű és hibrid felügyelt tanuló modellek performancia metrikái iparági megoszlásban	159
7. sz. melléklet:	Felügyelt egyszerű és hibrid módszerek performancia metrikái szűkített és a teljes adatvagyonon	160
8. sz. melléklet:	Egyszerű és hibrid felügyelt módszerek ROC és PR-görbéi, valamint tanulási görbéi szűkített adatvagyonon (könyvvizsgálathoz kapcsolódó informatikai vizsgálat).....	161
9. sz. melléklet:	Egyszerű és hibrid felügyelt módszerek ROC-görbéi és tanulási görbéi szűkített adatvagyonon (pénzügyi szektor).....	163
10. sz. melléklet:	Felügyelt modellek kategóriáinak értékelése varianciaelemzéssel	165
11. sz. melléklet:	A genetikai potenciál kereső és modell-preferencia levezetésre használt eljárások pszudokódjai	166
12. sz. melléklet:	A genetikai potenciál kereséshez alkalmazott alappopuláció nyersadatai	171
13. sz. melléklet:	A genetikai potenciál kereséshez alkalmazott alappopuláció rangsorai attribútumonként	172
14. sz. melléklet:	A genetikai potenciál kereséshez alkalmazott populáció rangsorai a második iterációban.....	173
15. sz. melléklet:	A genetikai potenciál kereséshez alkalmazott populáció nyersadatai a 11. iterációban.....	174
16. sz. melléklet:	A genetikai potenciál kereséshez alkalmazott populáció rangsorszámai a 11. iterációban.....	175
17. sz. melléklet:	Felügyelet nélküli modellek objektumleíró tulajdonságai	176
18. sz. melléklet:	Felügyelet nélküli modellek objektumleíró tulajdonságainak irány-preferenciát ellenőrző korrelációs mátrixa.....	179
19. sz. melléklet:	Felügyelet nélküli modellek leíró statisztikái.....	180
20. sz. melléklet:	Felügyelet nélküli modellek szórás irány-preferenciáinak értékelése varianciaelemzéssel.....	181
21. sz. melléklet:	Modell-leíró tulajdonságokra illesztett anti-diszkriminatív függvény súlyszámai (első irány-preferencia nézet).....	181
22. sz. melléklet:	Felügyelet nélküli modellek leíró tulajdonságainak értékelése varianciaelemzéssel (első irány-preferencia nézet).....	182
23. sz. melléklet:	Modell-leíró tulajdonságokra illesztett anti-diszkriminatív függvény súlyszámai (második irány-preferencia nézet)	183
24. sz. melléklet:	Felügyelet nélküli modellek leíró tulajdonságainak értékelése varianciaelemzéssel (második irány-preferencia nézet)	184
25. sz. melléklet:	Felügyelet nélküli modellek kategóriáinak értékelése varianciaelemzéssel	185

Jelölések jegyzéke

$\mathbf{x}$	bemeneti vektor
$\mathbf{X}$	bemeneti mátrix
$\mathbf{f}$	bemeneti mátrix attribútum-vektora
$\mathbf{y}$	tényadat (célváltozó)
$\hat{\mathbf{y}}$	becsült adat
$\mathbf{w}$	súlyvektor
$\mathbf{c}$	osztályozó, ahol $\mathbf{y} = \mathbf{f}(\mathbf{x})$ a tökéletes matematikai leképezés, így $\mathbf{c} = \hat{\mathbf{f}}(\mathbf{x})$
$\mathbf{C}$	együttes osztályozó
E	kalkulált hiba pl. átlagos négyzetes hiba
α	koefficiens
γ	tanulási ráta
$\mathbf{g}$	gradiens
$T(\mathbf{X}, \mathbf{y})$	tanulási függvény
$P(\mathbf{c}, \mathbf{X})$	predikciós függvény
$L(\hat{\mathbf{y}}, \mathbf{y})$	veszteségfüggvény

1. BEVEZETÉS

1.1. A téma aktualitása, célkitűzései

Folyamatos és növekvő igény mutatkozik az üzleti folyamatok automatizálására (STEGMAN et al. 2019), mely ma már nem kizárólag a rutinmunkának tekinthető, vagyis pl. szabályalapú tevékenységek kiváltására korlátozódik, hanem a komplexebb, kreativitást és akár szubjektív értékítéletet is igénylő feladatok gépiesítésére is kiterjed. A jelen kor technológiái, például a Big Data, Mesterséges Intelligencia, Felhő-alapú Megoldások, valamint a hardveres erőforrások gyarapodása (pl. HPC – High Performance Computing) és teljesítményük fokozatos javulása lehetővé teszik a magas színvonalú automatizált döntéselőkészítést és munkavégzést, mely szakmai és kutatási terület évről évre egyre bízhatóbb eredményeket tudhat magáénak (pl. BODA 2019, OLÁH et al. 2019, LENCSÉS et al. 2019).

Mindazonáltal, az automatizálást támogató, integrált, üzleti-informatikai megoldások fejlesztése és az eddigi elavult rendszerek korszerűsítése időről időre kihívások elé állítják a technikai szakembereket, vezetőket és felhasználókat, mivel olyan kockázatokat (pl. a kibertérben tárolt üzleti adatvagyon bizalmasságának sérülése) hordoznak, melyek szükségessé teszik az információ kezelésével, biztonságtechnikájával és ellenőrzésével foglalkozó funkcionális egységek megjelenését - megteremtve az informatikai erőforrásokra irányuló kockázatmenedzsment folyamatot (BARTA – GÖRCSI 2021).

A kockázatelemzési eljárások végső kimenete egy stratégiai szintű döntés a szervezetben felmerülő informatikai kockázatok csökkentésére (VASVÁRI 2008). Az informatikai kockázatelemzés eredményére támaszkodó döntés lehet a felmerülő kockázatok enyhítésére irányuló intézkedések megtétele, azaz a megismert kockázatokat kontrollfolyamatokkal történő mérséklése, úgy, mint a megelőző óvintézkedések, korrektív- vagy felderítő tevékenységek bevezetése.

Az implementált informatikai kontrollok meglepte még nem szavatolja a hibamentes működést, vagyis a vezetésnek rendszeresen meg kell győződnie arról, hogy a mitigáló intézkedések hatékonyan működnek, nem lehetséges azok megkerülése, illetve akaratlagos kijátszása (POMPON 2016). Erre egy gyakorlati példa az alkalmazások jogosultságkezelését megkerülendő hibajavításra (éles programfejlesztésre) szolgáló jogokkal történő visszaélés, melyekkel az alkalmazás szintű biztonsági logikát felül lehet írni, mely, ha a naplófájlok írási jogkörrel való hozzáféréssel párosul, akkor egy rosszindulatú felhasználó vagy támadó képes csalások kivitelezésére a rendszerben, úgy, hogy saját digitális lábnyomait törölni képes.

A belső informatikai kontrollok hatékonyságának feltárására szolgáló funkcionális terület megnevezése az informatikai audit-csoport, melyet olyan üzleti és informatikai szakmai tudással rendelkező szakemberek alkotnak, akik elsődleges célja a belső informatikai kontrollkörnyezet ellenőrzése és folyamatos tesztelése, annak garantálására, hogy az informatikai folyamatok teljes mértékben a szervezetek üzleti célkitűzéseit követik a lehető legnagyobb információbiztonsági szint fenntartása mellett. Az auditálás megjelenése és annak bevezetése egy szervezetben segíthet a kontrollhiányosságok feltárásában, azonban továbbra is fennáll az emberi tényező, azaz a véletlenül és/vagy szándékosan elkövetett hiba lehetősége, valamint az emberi erőforráskorlátok (BARTA 2018e). A legtöbb esetben szinte lehetetlen manuálisan végrehajtani az ellenőrzést annak dinamikus karakterisztikái végett, pl. azt, hogy a felhasználók kiosztott jogosultságai megfelelőek-e, úgy, hogy mindeközben a csalások elkövetésének kockázata minimális legyen a származtatott jogosultságokon keresztül, vagy, hogy az összes informatikai rendszer biztonsági konfigurációja kövesse az iparági gyakorlatokat.

Mindezért az auditorok egy előre meghatározott statisztikai mintaelemszámmal igyekeznek kiszűrni a nem-megfeleléseket, azaz az ellenőrzés csak részleges (BARTA 2018e). Továbbá, számos példa tetten érhető hazánkban és külföldön is, amikor az auditori munka minősége megkérdőjelezhető volt függetlenségi konfliktus, szervezeti érdekellentétek vagy az etikátlan üzleti magatartás következtében:

- Gondoljunk csak a 2001-es Enron botránya, melynek következtében a világ 5. legnagyobb auditor cége, az Arthur Andersen megszűnt, mivel az Egyesült Államok Legfelsőbb Bírósága felelősnek ítélte szakmailag és erkölcsileg is kifogásolható gyakorlatai miatt (GREENHOUSE 2005).
- Magyarországon is említhető negatív példa, ami kapcsolatba hozható az informatikai ellenőrzések hanyagságával:
 - 2015. február 24-én a Magyar Nemzeti Bank (MNB) felfüggesztette a Buda-Cash Bróker Zrt. működését (Magyar Nemzeti Bank 2015),
 - amely pár hétre rá magával sodorta a Hungária Értékpapír Zrt.-t és
 - Quaestor Értékpapír-kereskedelmi és Befektetési Zrt.-t, melyek együttvéve megközelítőleg 250 milliárd forintnyi fiktív kötvényt bocsájtottak piacra. Lukács János, a Magyar Könyvvizsgálói Kamara elnöke, a 2015. április 8-án adott interjújában kiemelte a mélyrehatóbb informatikai ellenőrzések fontosságát és igényét a jövőbeli botrányok elkerülése végett, és a könyvvizsgálók szerepét az ehhez hasonló események megelőzése érdekében (Magyar Könyvvizsgálói Kamara 2015). A Budacash és Questor cégek esetében az informatikai rendszerekben tárolt adatok kerültek meghamisításra, melynek révén történtek a csalások. Látszólag sem a belső, sem a külső fél nem használt alaposan megtervezett audit-eljárásokat, mivel adott volt a lehetőség a fiktív kötvények kibocsájtására, amit a szervezetek információrendszerei valósnak kezeltek. Ez az eset, mind az alaposabb informatikai ellenőrzés bevonását és teljeskörű minőségbiztosítását, mind újabb jogszabályi megfeleléség szerepét újraértékeli.

A felsorolt negatív példákból és az audithoz köthető kockázatokból egyértelművé válik, hogy a téma aktuális, és ez a közeljövőben még inkább hangsúlyos és jelentős lesz (vö. Ipar 4.0, MI-koalíció), mivel az informatikai iparosítás és üzleti automatizáció által nyerhető üzleti előnyök még inkább késztetni fogják a versenyben résztvevőket a technológiai megoldások implementálására (VINOGRADOV 2020).

Az auditori munka hatékonyságát növelendő, olyan automatizált megoldás szükséges, mely képes objektíven a hiányosságokra utaló jelek felkutatására (üzleti probléma), feldolgozására és értelmezésére, nem kizárólag szabályalapú értékelésére, hanem a komplex, háttérben meghúzódó összefüggések feltárására is, melyet a klasszikus statisztikai módszerekkel nem lehet hatékonyan kezelni. Ezen felsorolt elvárások gyakorlati kikényszerítése szakirodalmi kutatásaim alapján a mesterséges intelligencia fogalmkörébe illeszthetők, így a mesterséges intelligencia, a levont következtetések szerint, létjogosultsággal rendelkezhet a probléma megoldását illetően mindenkor a Knuth-i (1995) elv követésével, ahol a Knuth-i elvárás látszólag egyszerű: „*Tudás/tudomány az, ami forráskódba átírható*”, mégis a XXI. század alapelvárásaként azonosítandó be (PITLIK et al. 2017).

A mesterséges intelligenciával támogatott eszközök fejlesztésének jelentőségét alátámasztandó, hazánkban is elkészült 2020 májusában Magyarország Mesterséges Intelligencia Stratégiája (2020) a Mesterséges Intelligencia Koalíció jóvoltából, mely dokumentum részleteiben kitér az

intelligens rendszerek bevezetésének és üzemeltetésének üzleti eseteire, célokat és akcióterveket fogalmaz meg a következő tíz évre, 2030-ig.

Az intelligenciával ellátott szoftveres megoldások pl. tanulási mechanizmusaik révén képesek lehetnek a gyanús tevékenységek monitorozásában, ezzel hatékony módon valós értéket képezve az informatikai auditban. A felvetett probléma, tehát, az informatikai rendszerekben elkövetett kontrollhiányosságok, csalások, anomáliák, tehát gyanús események, mint informatikai kockázat fogalmának döntési helyzet-specifikus kialakítása mesterséges intelligencia módszerek alkalmazásával, vagyis az élő emberi képességek számítógépes leképezésével, helyettesítésével, támogatásával. Kiemelendő, hogy a dolgozatban prezentált kutatás és eredmények sem képesek az előző bekezdésben tárgyalt masszív csalások ellen védelmet biztosítani, s így a dolgozat célja sem az, hogy tetszőleges hatáskörrel bíró bűnözők elleni megoldásokat kínáljon. A kutatás elsődleges kiindulópontja, hogy egy cégvezetés legális körülmények között a biztonsági környezet hatékonyabb üzemeltetése céljából, minőségbiztosítási szempontokat előtérbe helyezve kíván megbizonyosodást szerezni a kontrollkörnyezet eredményességéről.

A mesterséges intelligenciával támogatott szoftveres megoldások egyik legnagyobb próbatétele a rendszerek tanítására és tesztelésére felhasznált minták korlátozott elérhetősége és annak diverzitása, így szükségszerű olyan módszerek kutatása, melyek képesek a rendelkezésre álló adathalmazban rejlő maximális potenciál (információs többletérték) kiaknázására. A dolgozat egyik centrális eredménye az elérhető adatok felhasználási optimumának letapogatása (genetikai potenciál keresése), egyrészt a tanuló minta hatékonyabb feldolgozásán, másrészt a klasszikus tesztelési eljárások nélküli modell-preferencia levezetésén keresztül.

Összefoglalóan, a kutatási cél olyan modellek megalkotása, melyek az informatikai rendszerekben felmerülő kontrollhiányosságok automatizált felderítésére (gyanúgenerálásra) irányulnak a mesterséges intelligencia alapú fogalomalkotás lehetőségeit felhasználva a modellalkotásban és a modelljóság mérésben egyaránt, ahol a rendelkezésre álló audit naplóállományból kell a mesterséges, elemi/mérhető és optimalizáltan aggregálható kockázatfogalom megalkotása keretében ennek mértékét, normáját algoritmikusan levezetni a tanításra alkalmazott adatok, lehetőleg optimális felhasználásával. Tehát a cél egy robot-auditor fejlesztése, mely alkalmas elrugaszkodni az emberi önkényességtől és az audit egyes részfeladatainak automatizált elvégzésére képes, és így a dolgozatban, minden olyan automatizált megoldást robotizációnak nevezek, mely el tud távolodni, akár minimálisan is a manuális feldolgozástól és emberi szubjektivitástól.

A kutatás célkitűzései (Ci) az alábbi pontok szerint strukturálhatók:

- **C1:** A Knuth-i elvet követve az információbiztonsági auditok hatékonyságát növelendő, létrehozandó olyan mesterséges intelligenciával ellátott döntéstámogató rendszer (robot-auditor), mely automatizáltan a historikus információbiztonsági auditjelentésekből tanulva képes a kontrollhiányosságok és kontrollterületek közötti összefüggések matematikai feltárására és javaslattevételeivel a potenciális emberi hibából fakadó észlelési kockázatok csökkentésére.
- **C2:** A fejlesztendő mesterséges intelligenciával ellátott szoftveres robot-auditornak kényszerűen alkalmasnak kell lennie a rendelkezésre álló adathalmaz minél inkább az optimálishoz közeli felhasználására, mely által teljesítménye maximalizálható, azaz a cél a robot-auditor genetikai potenciáljának kiaknázása a tanulási adathalmaz irányított redukálása révén.

- **C3:** A tanulásra felhasznált adathalmaz információtartalmát növelendő, anti-diszkriminatív módon szükséges az egyes robot-auditor alternatívák teljesítményeinek összehasonlítása, a legjobb alternatíva kiválasztása, melyhez nem szükséges validációs és teszhalmaz elkülönítése a szokásos tesztelés általi adat/információ-vesztési gyakorlattal szemben.

A célok megfogalmazásához a SMART kritériumrendszert rendelttem, melynek megfelelosége biztosítja, hogy a kutatási célkitűzések mennyiségi és minőségi feltételei teljesíthetők. A SMART feltételrendszer az alábbi elemeket tartalmazza YEMM (2012) alapján:

- Tényleges (Specific)
- Mérhető (Measurable)
- Teljesíthető (Achievable)
- Releváns (Relevant)
- Időhöz kötött (Time-based)

Az 1. táblázat szemlélteti a célkitűzések kiértékelését a SMART kritériumok mentén.

1. táblázat: A kutatás célkitűzéseinek értékelése a SMART feltételrendszer alapján

Kritériumok	C1	C2	C3
<i>Tényleges</i>	Konkrét üzleti probléma megoldására vonatkozik	Módszertani megoldások és gyakorlatok fejlesztésére vonatkozik	Módszertani megoldások és gyakorlatok fejlesztésére vonatkozik
<i>Mérhető</i>	A cél teljesülése objektíven visszaellenőrizhető	A cél teljesülése objektíven visszaellenőrizhető	A cél teljesülése objektíven visszaellenőrizhető
<i>Teljesíthető</i>	A cél megvalósításához szükséges adatok primer forrásból beszerezhetők, a modellezés a rendelkezésre álló eszközök által kivitelezhető	A modellezéshez szükséges hardveres kapacitások elérhetők	A modellezéshez szükséges hardveres kapacitások elérhetők
<i>Releváns</i>	Időben aktuális üzleti probléma megoldására vonatkozik	Időben aktuális módszertani eljárások fejlesztésére vonatkozik	Időben aktuális módszertani eljárások fejlesztésére vonatkozik
<i>Időhöz kötött</i>	A rendelkezésre álló idő alatt tervezetten teljesíthető	A rendelkezésre álló idő alatt tervezetten teljesíthető	A rendelkezésre álló idő alatt tervezetten teljesíthető

Forrás: Saját szerkesztés

1.2. A dolgozat szerkezete

A kutatás keretmodelljét és teljes folyamatát, tehát a célkitűzések elérésének útját az 1. számú melléklet szemlélteti, mely összefoglaló ábraként és további referenciaként funkcionál az alkalmazott modellezési gyakorlatok és műszaki lépések rendszerezésére. A dolgozat 10 fejezete ezen kutatási folyamat szakaszait és azok eredménytermékeit ismerteti:

1. BEVEZETÉS: A téma felvezetése, aktualitásának alátámasztása, a célkitűzések ismertetése. A kutatási célok meghatározásának elsődleges szempontja a tudományos társadalom és az információs többletértéket valóban realizálni képes potenciális célcsoportok részére történő magas színvonalú és hozzáadott értéket képviselő kutatási eredmények publikálása volt felhasználva a jelenkor technológiai adottságait, mely középpontjában az auditok automatizált minőségbiztosítására irányuló üzleti probléma állt.

2. SZAKIRODALMI ÁTTEKINTÉS: A szakirodalmi áttekintés a deklarált célok meg nem oldottságának bizonyítását, s így a kutatás szükségességét, annak racionalitását kívánja alátámasztani, melyek kifejtésre kerülnek az egyes alfejezetek összefoglalásában. A mindenkori cél az objektív megalapozottság, a logikus felépítés és ok-okozati összefüggések feltárása. A kutatás alapköveként szolgáló szakirodalmi kutatás fókuszja a kutatás témájául választott információbiztonsági auditok hatékonyságának növelését célzó módszertanok megismerése, gyenge pontjainak feltárása, és azok feloldását támogató megoldások potenciális továbbfejlesztése, mely során a szakmai követelmények azonosítása és szakmai tartalomban való elmélyülése kitüntetett szereppel bír. A szakirodalom egyszerre mélyreható és széles spektrumú feltérképezése és kritikai elemzése segítséget nyújtott a kutatási célokhoz tartozó hipotézisek finomhangolásához és megerősítette azok racionális mibenlétét.

Az első alfejezetben **(2.1. Az interdiszciplináris kutatás tudományterületi kapcsolatrendszer)** a kutatás tudományterületi elhelyezését kísérel meg, mely egyértelművé teszi annak diszciplináris osztályozását és a kutatás alapelveit, valamint, hogy a mesterséges intelligencia esetében minden megközelítés kényszerűen interdiszciplináris illik, hogy legyen.

A második alfejezet **(2.2. Információrendszerek, biztonság és audit)** az információ, információrendszerek és azok biztonságtechnikájára és auditálására vonatkozó definíciókat, tudományos és szakmai nézeteket részletez. Bemutatja a kutatás üzleti területét és kihívásait, alappillérként szolgál az azonosított üzleti probléma sajátosságainak szemléltetéséhez és megértéséhez, meghatározza annak kereteit, határait és kockázatait.

A harmadik alfejezet **(2.3. Mesterséges intelligenciával ellátott modellfejlesztés)** a kutatási problémára irányuló módszertan, azaz a mesterséges intelligenciával ellátott szoftverfejlesztéssel szemben támasztott követelmények és elvárások fogalmi meghatározásait tárgyalja, ismerteti a kutatás szempontjából relevánsnak és/vagy kritikusnak tekinthető megoldások karakterisztikáit, rendszerezi az elméleti megközelítéseket.

A negyedik alfejezet **(2.4. Gyanúgenerálás az információbiztonság kutatási területén)** a gyanúgenerálás értelmezését és a kutatással szorosan kapcsolatban álló azokat erősítő, vagy ellenpontosító szakirodalmi kutatási eredményeket hivatott bemutatni.

Az ötödik alfejezetben **(2.5. Hipotézisek felállítása a szakirodalmi áttekintés alapján)** kerülnek ismertetésre a szakirodalom elemzése és kiértékelése által racionálisnak vélt hipotézisek, melyek összhangban állnak a kutatási célkitűzésekkel és alátámasztásuk a soron következő fejezetek tárgyát képezi.

3. ANYAG ÉS MÓDSZER: A fejezet ismerteti a kutatás során alkalmazott adatgyűjtés technikáját, az adatbázisra vonatkozó minőségi követelményeket, a modellezési gyakorlatokat és alkalmazott jószágmetrikákat. A dolgozatban bemutatott kutatás alapos tervezési munkát igényelt, ezért annak folyamata előre meghatározott, részletes szakmai és tudományos módszertani

irányelveknek megfelelően lett strukturálva a mindenkor célkitűzéseket fókuszban tartva. A tervezés stratégiai szintjén a legfontosabb az volt, amit minden mesterséges intelligencia-alapú projekt/kutatás esetén minimum-követelményként kell betartani: a „jó” fogalmát kell előre és minél pontosabban algoritmizáltan (Knuth-i alapon) definiálni.

Az első alfejezet (**3.1. Adatgyűjtés**) részletezi a terepmunka során gyűjtött primer adatok forrását, mennyiségét és minőségét, valamint ismertetésre kerülnek az adatbázis struktúrája és korlátai.

A második alfejezet (**3.2. Alkalmazott algoritmusok és statisztikai eljárások**) szolgáltatja a hipotézisek bizonyításához felhasznált modellezési gyakorlatokat és technikákat, melyek bemenetét a terepmunkán gyűjtött adatvagyon, valamint kapcsolódó számítási eredmények képeztek.

A harmadik alfejezet (**3.3. Gépi tanuló rendszerek kiértékelése**) a kutatás objektív eredményeinek megállapítására és a modellek jóságának mérésére szolgáló metrikákat prezentálja.

A negyedik alfejezet (**3.4. Felhasznált eszközök és technológiai megoldások**) a dolgozatban közölt kutatás elkészítéséhez és modellezés megvalósításához felhasznált eszközöket ismerteti.

Az ötödik alfejezet (**3.5. A kutatási célok és hipotézisek rendszere**) összefoglaló jelleggel taglalja a kutatási célkitűzések és hipotézisek rendszerét, valamint az alkalmazott módszereket.

4. EREDMÉNYEK: A fejezet részletezi a felállított hipotézisek bizonyítását a terepmunkán gyűjtött adatok elemzésére alapozva a 3. fejezetben ismertetett matematikai apparátusok felhasználásával.

Az első alfejezet (**4.1. A kutatás során gyűjtött adatok leíró statisztikái**) a kutatás során gyűjtött adatvagyon leíró statisztikáit mutatja be, mely ezáltal magas szinten taglalja a primer kutatás során gyűjtött és rögzített adatvagyon karakterisztikáit.

A második alfejezet (**4.2. Gyanúgenerálás információbiztonsági kontrollhiányosságok detektálására**) az első célkitűzés elérését, és ahhoz kapcsolódó hipotézisek bizonyítását szolgáltatja, mely különböző, a mesterséges intelligencia világában alkalmazott modellek primer adaton történő futási eredményét és teljesítményét ismerteti.

A harmadik alfejezet (**4.3. Genetikai potenciál keresése a gépi tanulás adatvagyonának redukált felhasználásával**) a második célkitűzés elérését, és ahhoz kapcsolódó hipotézis bizonyítását részletezi, egy saját fejlesztésű kereső eljárás mechanikáját és működését szemlélteti elméleti alapokon és gyakorlati példával demonstrálva a primer kutatás során gyűjtött adatok bevonásával.

A negyedik alfejezet (**4.4. Modell-preferencia levezetése klasszikus tesztelési eljárások nélkül**) a harmadik célkitűzés elérését, és ahhoz kapcsolódó hipotézis bizonyítását szolgáltatja, véletlen modellek között a legideálisabb modell levezetését prezentálja saját fejlesztésű innovatív algoritmus alkalmazása által.

Az ötödik alfejezet (**4.5. A dolgozat célkitűzéseinek teljesítése a SMART feltételrendszer alapján**) a kutatás célkitűzéseinek teljesülését értékeli.

5. ÚJ ÉS ÚJSZERŰ TUDOMÁNYOS EREDMÉNYEK: A fejezetben kerülnek kifejtésre a modellezési gyakorlatok kiértékelése alapján megállapított innovatív, új és újszerű tudományos eredmények, mely összefoglaló jelleggel ismerteti a dolgozat az emberiség tudásbázisához tett hozzáadott értékét.

6. KÖVETKEZTETÉSEK ÉS JAVASLATOK: A fejezet részletezi a kutatás és kísérletek által nyert gyakorlati alkalmazhatóság területeit, egyfajta jövőkép jelleggel összefoglalja azon javasolt kutatási irányokat, mellyel a dolgozatban közölt eredmények javíthatók és alapot szolgáltat a további kutatási célkitűzések meghatározásához, valamint gondolatokat fogalmaz meg a kutatási eredmények piacosítható jellegéről.

7. ÖSSZEFOGLALÁS: A dolgozat, kutatási célok és eredmények összefoglalása.

8. SUMMARY: A dolgozat, kutatási célok és eredmények összefoglalása angol nyelven.

9. IRODALOMJEGYZÉK: A kutatás során felhasznált szakirodalom jegyzéke.

10. MELLÉKLETEK: A kutatáshoz kapcsolódó egyéb adatfeldolgozási eredményeket, mellékszámításokat, ábrákat, stb. tartalmazó fejezet.

A dolgozat az alábbi tartalmi és formai szempontok betartásával készült:

- A dolgozat nem csak szövegesen, hanem folyamatábrák és pszeudokódok felkínálásával is támogatni kívánja a magas komplexitású jelenségek, az egyes alrendszerek egymással való kapcsolatának megértését;
- A dolgozat *dőlt* betűvel szedetten tételes idézeteket tartalmaz. Minden más a szerző saját véleménye;
- A dolgozat él a kiemelés lehetőségével **vastagon** szedett kulcsszavak formájában a gyorsabb áttekintés támogatására;
- A dolgozat keretében felmerülő számértékek esetén csak a szükséges mennyiségű tizedesjegy (általában kettő) szerepel. A dolgozatban a tizedesjel a pont. A dolgozat nem tartalmaz ezres-határoló jelet, mert a kezelendő számok nagyságrendje ezt nem követeli meg;
- A dolgozat saját jelölésrendszert alkalmaz a szöveges részben és a pszeudokódok esetén.

2. SZAKIRODALMI ÁTTEKINTÉS

Jelen fejezet a kutatás szükségességét és annak racionalitását kívánja alátámasztani. A mindenkor cél az objektív megalapozottság, a logikus felépítés és ok-okozati összefüggések minél komplexebb feltárása, így a releváns szakirodalom elemzése és konstruktív kritikai értékelése építőként szolgál a probléma bemutatásának, megoldási megközelítéseinek és az eredmények minősítésének szempontjából.

2.1. Az interdiszciplináris kutatás tudományterületi kapcsolatrendszere

A kutatás alapvető célkitűzése sok dimenzió mentén, azaz a mindenkor vizsgált objektumok (folyamatok, rendszerek, auditok, kontrollok, stb.) rendelkezésre álló leíró attribútumainak konstellációjából adódó maximálist hatékonyan közelítő tudás kinyerésével a kontrollhiányosságok detektálása, ahol a rendszer által kimenetként szolgáltatott célváltozókra adott becslés értéke a lehető legjobb, vagyis az aggregált jószágmetrikák értéke optimális/fenntartható szemben a nem aggregált és nem optimalizált megoldások egysíkúságából fakadó dinamikus kockázatokkal. A minél inkább automatizálható és minél inkább valós idejűséget közelítő és biztosító rendszernek képesnek kell lennie, akár egy előre még általa nem ismert auditról és kontrollról objektív véleményt alkotva meghatározni, hogy az milyen mértékben gyanús, azaz valószínűsíthető a kontroll hatékonyságának kifogásolhatósága. Egy ilyen alkalmazás elsődleges célközönsége az audit vezetője, aki az audit típusából, hatóköréből, az auditált szervezet ismeretében (korábbi tapasztalatok, iparági sajátosságok, működési környezet) képes akár az audit előtt, annak tervezési fázisában egy előrejelzést szerezni a magas kockázatú kontrollterületekről, vagy utólagos ellenőrzésként megbizonyosodni arról, hogy az audit sikeresen feltárta az összes kontrollhiányosságot. Az audit vezetője a rendszer kimenete által hasznos bemenetre tehet szert, mivel az audit tervet a potenciális hiányosságok meglétének valószínűsége fókuszáltabbá teheti, így döntéstámogató funkcióval bír. Az alkalmazás célcsoportjaiba, továbbá, tartozhatnak különböző szakmai területi vezetők (pl. informatikai vezető, információbiztonsági vezető, kockázatkezelésért és működésért felelős vezető, jogi vezető, stb.), akik a szakterületük által működtetett kontrollok megfelelőségét képesek adatvezérelten ellenőrizni, a gyanús eseteket kockázat alapon célozottabban megvizsgálni, evégett, az alkalmazás egy döntéselőkészítő folyamat keretrendszerébe illeszthető.

Kijelenthető, hogy a dolgozatban bemutatásra került kutatás a döntéstámogatás tudományos elméleteit felhasználva, alkalmazva és meghaladva, a döntéshozatali folyamat hatékonyságát hivatott növelni, a döntéstámogatás minőségi kereteit és elvárásait szándékozik a mesterséges intelligencia által biztosított szemléletmódban javítani, így annak tudományos kategorizálását illetően a célkitűzés szemszögéből vizsgálva, a kutatási terület a döntéselmélet nómenklatúrájába sorolható. Az elsődleges célkitűzés eléréséhez felhasznált eszközrendszer fejlesztésére irányuló kutatási tevékenységek és megközelítések irányából szemlélve, tehát módszertani aspektusban értelmezve, a kutatási terület a számítástechnika tudományágaként osztályozható, ami a döntéselmélet modernkori segédtudományának tekinthető (ZOLTAYNÉ 2002).

A Magyar Tudományos Akadémia (2017) Tudományági nómenklatúrája alapján a kutatás elsősorban a IX. Osztály. *Gazdaság- és jogtudományok* alágának *Gazdaságtudomány*, azon belül a *Gazdasági operációkutatás és döntéselmélet* osztályába sorolandó.

A kutatás metodológiáját tekintve pozitivista szemléletű, azaz a hangsúly a létező valóság minél inkább objektív megismerésén és felderítésén alapszik, ahol az objektivitás egy triviális értelmezése az előre levezetett eredmények jövőbeni visszaigazolódásának minél magasabb

aránya. A normativista paradigmával ellentétben, ahol a cél a szubjektív értékítéletek mentén történő vizsgálat, a pozitivizmus képes a társadalomtudományi kutatást is a természettudományok kutatási megközelítéséhez közelebb hozni. Ez úgy érhető el, hogy a kutatásban alkalmazott módszertanok kikényszerítik az adott társadalomtudományi kutatás hatóköre alatt gyűjtött, a szubjektivitást megengedő adathalmaz szilárd/konzisztens matematikai apparátusok szerint történő feldolgozását (KÁSA 2011). Mivel a társadalomtudományi mérések esetén nincs fizikai mérőeszköz, ami szavatolná a teljes objektivitást, ezért a pozitivizmus által megkövetelt elvárások részben sérülnek, tehát a kényszerű normativizmus külső adottság. A kutatónak szükséges minden olyan módszertani eszközt alkalmaznia, amivel képes a szubjektivitásból adódó pontatlanság kockázatát minimalizálni. Véleményem szerint csak ezzel a megközelítéssel határolható el az emberi belemagyarázás veszélye a kutatási eredményekből, ezért a dolgozat szigorúan a pozitivisták paradigmát szándékozik követni, mely COMTE (2009) értelmezésében:

- „*valóság az elképzelttel szemben*”;
- „*haszon a haszontalansággal szemben*”;
- „*bizonyosság a határozatlansággal szemben*”;
- „*pontosság a bizonytalansággal szemben*”;
- „*viszonylagos az abszolúttal szemben*”;
- „*pozitívum a negatívummal szemben*”.

A fentiek alapjaiban felelnek meg a Knuth-i elvárásnak is, hiszen az emberi belemagyarázó készség területén tűnik egyedül feleslegesen robotizálni az emberi képességet. A robot számára az adatvezéreltség az, ami hasznosan képes az emberi intuíciót támogatni. Az embernél emberibb robotra nincs szükség vélhetően.

2.2. Információrendszerek, biztonság és audit

2.2.1. Az információ fogalmi megközelítései

Az információt egyes szakirodalmak az 5. termelési tényezőként tartják számon, többek között BODA et al. (2009), KISS (2016) és FARKASNÉ és MOLNÁR (2017), akik tanulmányukban az információt a tudással is azonosítják, mely ezen kutatás szempontjából máshogy értelmezendő, ugyanis a dolgozatban adat és információ minden, ami jelkét feldolgozásra kerül az algoritmusok által, melyek a tudás hordozói. POÓR et al. (2020) az információt és tudást külön fogalomként értelmezi. FARKASNÉ és MOLNÁR (2017) egy sajátos termelési tényezőként kezeli az információt (a tudással együtt), melyet a tudományos kutatás termel, és hozzájárul a teljes társadalmi átalakuláshoz, mivel az áthatja a gazdasági folyamatok egészét. Hasonló véleményen van MÁTYUS (2015), aki hozzáteszi, hogy a technológiai fejlődések révén a gyors információáramlás és feldolgozás, valamint az információból kinyerhető tudás létrehozta az információs társadalmat.

HARKEVICS (1960) az információ értékéről beszél: csak akkor válik az információ értékké, amennyiben hozzájárul egy kitűzött cél megvalósításához. Az információ közvetítése, értelmezése és feldolgozása nagyban hozzájárul a vállalati döntéstámogatáshoz, annak eszköze, nélkülözhetetlen kiszolgáló eleme. Szemléltetésül, egy új termék bevezetéséhez piaci információra van szükség a versenytársakról, fogyasztókról, szabályozói környezetről, ajánlásokról, technológiáról, tehát bizonytalanságot csökkentő ismeretről beszélhetünk (CHIKÁN 2008).

ÁGOSTON és SZLUKA (1989) az információt a „*hatalom*” szóval kapcsolják össze, implikálva, hogy az információ előnyt jelent annak birtokosa számára. CAPURRO (1992) kiemeli, hogy az

információ elveszíti a kapcsolatot az emberi világgal, arra utalva, hogy nem kizárólag ember és ember között létezik információáramlás, hanem ember és gép, illetve gép és gép között is zajlik információ csere, azaz „informálódás”. FÜLÖP (1996) az ipari termelés ki nem merülő tartalékként kezeli. MUNK (2007) az információt a valóság visszatükröződéseként értelmezi.

A szakirodalomban az előző bekezdésekben megfogalmazottakat kiértékelve, véleményem szerint az információ a termelési tényezők egyik meghatározó eleme. A döntéstámogatáson keresztül beépül a szervezeti stratégia készítésébe, annak kivitelezésébe és a teljes stratégiai irányításba, ezért is kezelhetjük termelőeszközként, mert az információ feldolgozása hozzájárul a stratégiai döntéshozatalon keresztül a javak előállításához. Az információ fogalmát, azonban, a különböző tudományos szakterületek eltérően definiálják. Alapvetően az információelmélet, mint matematikai és hírközlési tudományterület, az információ feldolgozásával és annak értelmezésével foglalkozik, és mint annak egyik jeles képviselője, és a tudományterület atyja SHANNON (1948) az információt az adó és vevő közötti valamilyen üzenet közléseként, eseményként nevezte meg. FORGÓ (2011) ezt azzal egészíti ki, hogy az információelmélet szerint, az információ a kommunikációs folyamat mennyiségi mértékegysége. KOMENCZI (2011) alapján, az információ egy az agyunkban meglévő elképzelés, ami magába foglalja a tudást, belátást, felismerést. WORTH és GROSS (1977) szociológiai szempontból közelítette meg az információ fogalmát, melyet társadalmi folyamatként írtak le, a jelek közleményként való észleléseként, melyből jelentésre lehet következtetni. A modernebb definíciók az információ és Big Data fogalmát gyakran összekötik, melyben az információt felruházzák olyan leíró elemekkel, mint annak mennyisége, változatossága, redundanciája, mivel a jelen technológia és eszközök lehetőséget nyújtanak az adatok és információ folyamatos tárolására, gyűjtésére és másolására (LIN et al. 2016).

A világháló mindennapi alkalmazásával, annak jelentős elterjedésével az információ megosztása másodpercek alatt zajlik. 2020 végével az Internet World Stats (2020) becslése szerint 4.9 milliárd internet felhasználó volt jelen, mely 1271%-kal több, mint a 2000-ben mért adat. A hardveres tároló kapacitások árának folyamatos csökkenése miatt egyre több adatot vagyunk képesek tárolni. 1 gigabite adat 1980-ban átlagosan 200 ezer dollárba, 2000-ben 10 dollárba, 2009-ben csupán 10 centbe, míg 2017 végén fél centbe került (KOMOROWSKI 2014, KLEIN 2017). A hardveres erőforrások árának efféle dramatikus csökkenése azt eredményezi, hogy a szervezetek képesek a keletkező információ permanens tárolására hosszú évekre visszamenőleg is.

2.2.2. Információrendszerek megjelenése és biztonsági kérdései

A szakirodalmi kutatásom és az ez idáig felsorakoztatott megfogalmazások alapján, álláspontom szerint az információt, a disszertáció kapcsán különös tekintettel a kockázatokról szóló becslésekre vonatkozóan, erőforrásként célszerű értelmezni a gazdálkodás és szervezéstudományok területén, mely alapjaiban járul hozzá hatékony feldolgozása és értelmezése által a szervezeti folyamatok eredményes működéséhez és a döntéshozatal elősegítéséhez (GÖRCSI – BARTA 2018, 2019). Termelési tényezőként való értelmezésében az információ, mint olyan egyedülálló tulajdonságokkal rendelkezik. A többi termelési tényezővel ellentétben, gyakorlatilag korlátlanul sokszorozható, az információ nem véges (nem fogy el), hasznosítása során annak állománya nem csökken, továbbá, lehetséges pillanatok alatt terjesztetni, korunk technológiai szintje lehetővé teszi annak egyszerű tárolását és mindenkor rendelkezésre állását. Az információ folyamatos raktározásából, azonban csak akkor teremthető érték, ha abból a szervezet képes tudásanyagot létrehozni és azt hatékonyan felhasználni (WARD 1998). Mindezek mellett az információ avulása

az összes termelési tényező között a legnagyobb – a kockázatokról szóló információ sem viselkedik ilyen tekintetben másként.

A megnövekedett információmennyiség kezelésének szükségessége életre hozta az információmenedzsment tudományát (DOBAY 1997), melynek központi eleme az információrendszer. Az információrendszerek kutatása az 50-esek évekre nyúlik vissza, melynek előfutárai a termeléstámogató-rendszerek voltak, majd a 70-es években megjelentek a vezetőket támogató információrendszerek, a 90-es évektől, pedig már az alapvető üzleti folyamatokat is információrendszerek szolgálták ki (WARD 1998). Jelenleg az információrendszerek képesek a teljes ellátási-lánc menedzsment funkcióit ellátni (MAGDA 2015), integrált vállalatirányítási rendszerek által a globálisan működő szervezetek teljes üzleti folyamatát központilag lefedni, beleértve a különböző elektronikus kereskedelmi és digitális üzletviteli funkciókat, hatékony kommunikációs csatornát kiépítve a szállítók, vevők és az állami szervek felé is (SZALAY 2009, BARTA – GÖRCSI 2017).

LANGEFORS (1973) szakirodalmi kutatásom alapján, a legelső volt, aki az információrendszer fogalmát definiálta. Korai művében a döntéstámogatás szerepét emelte ki az információrendszereknek, mely publikáció óta az alapos fejlődésen ment keresztül. K. C. LAUDON és J. P. LAUDON (1991) már kiemelte az adatokra vonatkozó keresés, tárolás, továbbítás funkcióját, mely technikai oldalról hangsúlyozta az információrendszert. K. C. LAUDON és J. P. LAUDON (2015) két és fél évtized elteltével megfogalmaz egy üzleti definíciót is, melyben szervezeti és menedzselési megoldásnak nevezik az információrendszereket. SZEPESNÉ (2011) úgy fogalmaz, hogy az információrendszer *„fő célja az információ-előállítása, vagyis olyan célorientált üzenetek létrehozása, amelyek a címzett számára újdonságot jelentenek, bizonytalanságot szüntetnek meg és feladataik, döntéseik teljesítésében segítséget nyújtanak.”* Az információrendszer, tehát képes az üzleti adatok összegyűjtésére, továbbítására és azok teljes életciklus menedzselésére.

Az információrendszerek a legtöbb szervezetben bizalmas, üzletileg kritikus információt kezelnek (GÖRCSI et al. 2019), evégett az üzemeltetett adatbázisok és alkalmazások sérthetlenségének és bizalmasságának biztosítása stratégiai prioritás kell, hogy legyen minden szervezet életében. Véleményem szerint ez azt jelenti, hogy a tradicionális üzleti operációt kiváltandó, kockázatalapú működésre indokolt a szervezeteknek átállni, azaz az üzemeltetési modell újragondolására és átalakítására van szükség.

A szabályozói környezet Magyarországon és az Európai Unióban is nagy hangsúlyt fektet az információbiztonsági előírások betartatására. Magyarországon (és hasonlóan az egész világon) kiemelendő, a pénzügyi szektor az egyik legnagyobb mértékben szabályozott ágazat, melyet hazánkban a Magyar Nemzeti Bank kötelező érvénnyel rendszeresen ellenőriz. Az előző években megjelent a személyes adatok védelmére vonatkozó részletes követelményrendszer a GDPR (General Data Protection Regulation – Általános Adatvédelmi Rendelet), mely egy új korszakot nyitott meg hatályba lépése óta az Európai Unió és Európai Gazdasági Térség területén működő szervezetek és magánszemélyek részére (BARTA 2018b, BARTA et al. 2020).

A biztonságos üzemeltetési környezet fenntartásáért és ellenőrzéséért az információbiztonsági szervezeti egység felelős, függetlenül az informatikai üzemeltetési osztálytól, egy egészséges szervezeti hierarchiában. Az információbiztonsági szervezet tevékenysége kiterjed az információrendszerek, a hálózat és egyéb fizikai és logikai eszközök védelmi intézkedéseinek meghatározására, kikényszerítésére és folyamatos nyomon követésére. A cél a biztonsági incidensek előfordulásának minimalizálása, mely hatással lehet a belső működésre és reputációs

kockázatot is vonhat maga után. Információbiztonságot érintő incidensek alapvetően három különböző forrásból érkehetnek az Information Security Forum (2014) csoportosítása alapján:

- külső, a szervezettől kívülálló csoportoktól, személyektől, melyek lehetnek kiberbűnözők, versenytársak ipari kémek, hobby hackerek, szélsőséges esetben terroristák;
- belső, a szervezet dolgozói, akik szándékosan vagy véletlen bizalmas információt osztanak meg a külvilággal vagy az információ ismeretére jogosulatlan munkatársakkal;
- természeti események, melyek az információ elérhetőségét és rendelkezésre állását veszélyeztetik elsődlegesen pl. heves esőzés, mely megkárosíthatja a szerverközpontot.

Az információbiztonsági szintet informatikai és szervezeti kontrollok implementálásával lehet növelni. A kontroll egy olyan szervezeti és/vagy technikai mechanizmus, melynek célja a hibamentes működés szavatolása (BARTA 2018e). A kontrollok összességét, kapcsolatát, egymásra épülésének architektúráját kontrollkörnyezetnek nevezzük.

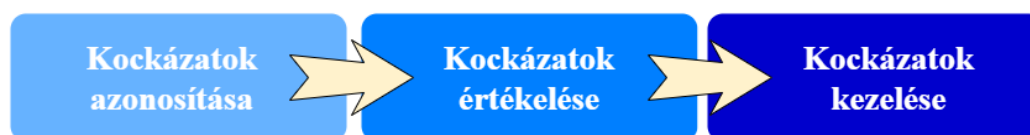
Alapvetően három különböző kontrolltípust különböztethetünk meg a káresemény kockázatának mitigálása szempontjából (POMPON 2016):

- preventív: egy fenyegetés bekövetkezési valószínűségét mérséklendő megelőző óvintézkedés pl. jelszavas védelem.
- detektív: egy már bekövetkezett káresemény felderítésére szolgáló utólagos mechanizmus pl. naplóelemzés.
- korrektív: egy már bekövetkezett káresemény kijavítását célzó tevékenység pl. adatok visszatöltése a mentési rendszerből adatvesztés esetén.

A kontrollok szükségességét, erősségét és rendszerességét a szervezeti információbiztonsági kockázatelemzés által lehet érvényre juttatni. A kockázatelemzés célkitűzése az erőforrások, fenyegetettségek és óvintézkedések (egy másik aspektusból vizsgálva: sérülékenységek) azonosítása, mely akkor eredményes, ha objektíven képes információval szolgáltatni a hiányosságokról (pl. a szervezet szerverszobája nem hatékonyan védett az illetéktelen hozzáférésektől, ezért preventív kontroll implementálása javasolt, mely lehet beléptetőrendszer telepítése, pin kód kikényszerítése, biometrikus azonosítás, vagy ezek együttes kombinációja) (BARTA - GÖRCSI 2020). Összefoglalva, az információbiztonsági kockázatelemzés segítséget nyújt azon területek feltárásában, melyek további biztonsági óvintézkedéseket igényelnek.

Az ISO (International Organization for Standardization – Nemzetközi Szabványügyi Szervezet) és az IEC (International Electrotechnical Commission – Nemzetközi Elektrotechnikai Bizottság) a 27000-es információbiztonságra vonatkozó szabványcsaládjában a következőként fogalmaz a kockázatról: „*A bizonytalanság hatása a célkitűzéseken*” (ISO/IEC 27000:2013(E) 2014). Ezt a hatást a szabvány pozitívként és negatívként is értelmezi, míg VASVÁRI (2018) a kockázatot kizárólag negatív hatású eseményként írja le, mely „*egy veszélyforrás képezte fenyegetés bekövetkezési lehetősége, amely kárkövetkezménnyel jár, és így kedvezőtlen hatást fejt ki az üzleti célokra.*” Az informatikai kockázat vagy kockázati tényező az ISACA (Information Systems Audit & Control Association – Információrendszer Audit és Kontroll Egyesület) IT kockázatkezelési keretrendszerének meghatározás alapján „*egy üzleti kockázat – azaz, az üzleti kockázat, mely szervezeti kereten belül kapcsolódik az információtechnológia felhasználásához, tulajdonlásához, működéséhez, bevonásához és befolyásolásához*” (YOUNG 2020).

Értelmezésem szerint az informatikai kockázatkezelés legvégső célja, hogy lehetőség szerint minél inkább automatizáltan, azaz a Knuth-i elvárásnak¹ megfelelően felszínre kerüljenek a szervezetben azok a hiányosságok, melyek az információ, mint szervezeti vagyon és termelési tényező, egyes attribútumainak, mint pl. teljesség, pontosság, megbízhatóság, elérhetőség, bizalmasság, hitelesség, stb. elvesztéséhez és sérüléséhez vezetnek. Ennek következtében, azon pontok és üzleti folyamatok, ahol a legkevesebb erőfeszítés összpontosul az információ biztonságos tárolásán, feldolgozásán és továbbításán. Az 1985-ben alapított COSO (Committee of Sponsoring Organizations of the Treadway Commission – A „Treadway Commission” Támogató Szervezeteinek Bizottsága) az egyik legnagyobb szervezetnek számít ma, mely szakmai publikációkkal és egy széles körben elismert keretrendszerrel támogatja a szervezeti kockázatmenedzsmentet, mely átfogja az információ, mint vállalati érték és termelési tényezők kockázatkezelését (COSO 2020). Az említett szervezet célja, hogy egy megalkotott modell biztosítása által (COSO modell) kockázatorientált megközelítésben bemutassa a szervezeti folyamatokat a vállalatok belső kontrollrendszere, irányítási rendszere, és a funkcionális területek tevékenysége alapján. Az ISACA hasonlóan több keretrendszert dolgozott ki. Érdemes kiemelni a CobIT-et (Control Objectives for Information and Related Technology – Információra és a Kapcsolatos Technológiára Vonatkozó Kontrollcélkitűzések), mely informatikai kontrollcélkitűzéseket határoz meg az üzleti kockázatok csökkentésére. 1996-ban publikálta első verzióját, melynek legfrissebb változata a CobIT 2019 címet viseli, és 2018 végén került kiadásra (HAES et al. 2018). Az ISACA másik széles körben alkalmazott publikációja a korábban említett IT kockázatkezelési keretrendszer, mely a nagyvállalatoknak nyújt segítséget az információbiztonsági kockázatmenedzsmenthez (YOUNG 2020). A COSO és Deloitte által kiadott „Kockázatelemzés a gyakorlatban” publikáció szerint a kockázatértékelés tevékenysége három főbb eljárásra bontható, melyet az 1. ábra szemléltet (CURTIS - CAREY 2012).



1. ábra: A kockázatelemzés magas szintű folyamata

Forrás: CURTIS - CAREY (2012)

- **Kockázatok azonosítása:** A folyamat részét alkotja szükségszerűen a kockázatelemzés hatókörének definiálása, a hatókörben lévő erőforrások üzleti hatáselemzése, a potenciális fenyegetettségek feltárása és bekövetkezési valószínűségek meghatározása.
- **Kockázatok értékelése:** A kockázatértékelés fázisában történik az értékelési kritériumok lefektetése és jóváhagyása, amely magában foglalja a kockázatok mértékének becslését, a nem várt hatásokat, a kockázatok előfordulásának gyakoriságát, azok prioritizálását.
- **Kockázatok kezelése:** A kezelési szakaszban a kockázati értékek hozzárendelése következik be a kockázati tényezőkhez a kritériumban definiáltaknak megfelelően, majd a kockázatkezelési terv elkészítése, mely célja az azonosított kockázatok mérséklését célzó korrektív akciótervek meghatározása.

¹ “Science is what we understand well enough to explain to a computer; art is everything else.” - „Tudomány az, amit értünk annyira, hogy elmagyarázzuk egy számítógépnek. Minden más művészet.” (KNUTH 1995)

YOUNG (2020) kiemeli, hogy nem kizárólag az egyes kockázatok egyéni hatásait szükséges figyelembe venni, hanem a kockázatok közötti interakciókat is, melyek hatásai együttesen még nagyobb kárt okozhatnak, azaz a kontrollok közötti kapcsolatrendszer kialakítása és ellenőrzése képes csak a teljeskörűség biztosítására. Egy kisméretű vállalat kevesebb figyelmet fordíthat az akár pénzügyi kimutatásokat (BARTA – LÉTEK 2015) is érintő ütemezett és automatizált eljárások futtatásának szabályozására, ami, ha párosul az összeegyeztethetetlen szerepkörök nem megfelelő szétválasztásával, akkor bekövetkezhet az az esemény, hogy egy egyszerű dialógus felhasználó változtatást indukál, és más időpontra helyezi a kritikusnak vélt feladatok elvégzését, azaz információbiztonsági incidenst idéz elő. Ez a példa felhívja arra a figyelmet, hogy nem elégséges elemi szinten vizsgálatot végezni, hanem a tényezők összességét szükséges minősíteni, tehát az összefüggések mértékét kell feltárni, mely komplex számításigényes feladat, és a probléma dinamikus jellegéből adódóan szabályalapú szakértői rendszerekkel kivitelezhetetlen. Álláspontom, hogy az elemi kockázatértékelés ezért nem szolgálhat minőségi hozzáadott értékkel.

A kockázatértékelési folyamat eredményterméke akkor érvényes, ha a kockázatelemzés képes volt összetett hatások értékelésére is, rangsorolva a kritikusabb területeket. A kockázati rangsornak, véleményem szerint, kényszerűen tükröznie kell az üzleti erőforrások (pl. informatikai alkalmazások), a releváns fenyegetettségek és sérülékenységek kollektív hatásait is. Az üzleti erőforrások és fenyegetettségek értéke adottság, ezért a kockázatok a sérülékenységek mértékének csökkentésével lehet mitigálni, mely az implementált kontrollok karakterisztikáitól függ. A kontrollok ellenállóképességét ellenőrizni szükséges, hogy a kockázatelemzés hiteles eredményeket szolgáltasson a vezetők számára, azaz ahhoz, hogy egy szervezet megbizonyosodjon, hogy az implementált kontrollok hatékonyan működnek, azokat auditáltatni kényszerül. Az audit funkció, mindezek értelmében, a szervezeti kockázatelemzés és kockázatkezelés szerves részét kell, hogy képezze.

2.2.3. Információbiztonsági kontrollok auditálása

Az információbiztonsági kontrollkörnyezet tervezet és implementációs szintű, valamint működési hatékonyságának ellenőrzésére és vizsgálatára specializálódó terület megnevezése az informatikai audit osztály (BARTA 2018e). Az auditorok olyan üzleti és informatikai szakmai tudással rendelkező szakemberek, akik elsődleges célja a belső és külső kontrollkörnyezet folyamatos tesztelése, annak meggyőződésére, hogy az implementált kontrollfolyamatok teljes mértékben a szervezetek üzleti célkitűzéseit követik a lehető legmagasabb információbiztonsági szint biztosítása mellett. Az informatikai auditok célja, lényegében, a szervezet vezetői és befektetői számára reális bizonyosságot nyújtani, hogy az üzleti folyamatokat támogató IT és információbiztonságot érintő kontrollok megfelelően (pl. követik az iparági jógyakorlatokat) lettek kialakítva, implementálva és hatékonyan funkcionálnak, valamint alkalmasak a biztonsági kockázatok elfogadható szintre való mérséklésére. A kontrollkörnyezetnek garanciát kell biztosítania MOLNÁR és KŐ (2009) alapján, hogy:

- az informatikai rendszerek helyesen és pontosan dolgozzák fel az üzleti adatokat (adatintegritás);
- az adatok hozzáférhetősége a legkisebb jogosultság elve alapján korlátozott, kizárólag annak címzettjei számára ismerhetők (bizalmasság);
- a belső eljárások szavatolják, hogy egy katasztrófahelyzet esetén is az adatok elérhetők legyenek, az üzletmenetfolytonosság biztosított (rendelkezésre állás);
- a szervezet munkatársai és partnerei betartják a biztonsági szabályzatokat, eljárásokat, irányelveket, melyeket a vezetés lefektetett, jóváhagyott és rendszeresen felülvizsgál.

Az audit osztálynak függetlenül kell operálnia minden funkcionális területtől és szükségszerűen közvetlen a felső vezetésnek kell jelentenie – az audit pártatlansága miatt (BUXBAUM 2006). Abban az esetben, ha egy szervezet nem engedheti meg magának a belső audit osztály fenntartását, lehetősége van azt kiszervezett szolgáltatásként igénybe venni. A külső tanácsadó auditorok a szervezet vezetésének jelentenek és hasonlóan funkcionálnak, mintha az egy a szervezeten belül függetlenül működő osztály szakemberei lennének. A szervezetek, azonban nem kizárólag belső bizonyosságszerzés miatt végeztethetnek auditokat, szükség lehet külső ellenőrzésre is jogszabályi megfelelés biztosítása érdekében pl. a pénzügyi szervezeteket Magyarországon az MNB rendszeres időközönként vizsgálja, emellett, számos szolgáltatás bevezetése esetén megköveteli a külső független audit elvégzését. Például a 26/2020. (VIII. 25.) MNB rendelet kötelezően előírja a kétéves független vizsgálatot Pmt. (2017. évi LIII. törvény a pénzmosás és a terrorizmus finanszírozása megelőzéséről és megakadályozásáról) szerinti távoli ügyfélazonosításra képes hírközlő eszköz alkalmazása esetén. A könyvvizsgálatra kötelezett szervezetek külső könyvvizsgáló általi ellenőrzésében is megjelenik az információbiztonsági kontrollok ellenőrzése, mivel a számviteli adatok sértetlenségének szavatolását negatívan befolyásolhatja a hiányosan működő, a csalás és téves könyvelési tételek kockázatának valószínűségét alacsony színvonalon mitigáló óvintézkedések is (BAESENS et al. 2015).

A kontrollkörnyezet hatékony működésének ellenőrzésére az audit osztály szakemberei teszteket végeznek egy előre meghatározott, a vezetőség és befektetők által felülvizsgált és jóváhagyott auditterv alapján. A legtöbb esetben a tesztelés statisztikai mintavételezési eljárással történik, melyben szerepet játszanak a tesztelésre kiválasztott kontrollok gyakoriságai. Például, belépő munkatársak jogosultságainak ellenőrzése egy multinacionális szervezet esetén lehet napi feladat, azonban az adatgazdák általi jogosultságfelülvizsgálat meglétének ellenőrzése általában csak évente egyszer történik. Az audit a kontrollok megfelelő üzemeltetéséről alátámasztó evidenciák (pl. naplózó rendszerből exportált konfiguráció) elemzése által bizonyosodik meg, mely szerves részét képezi az audit dokumentációjának. Kiemelendő, az audit sem tud abszolút bizonyosságot szerezni a kontrollok kifogástalan működéséről, mivel kvázi lehetetlen az összes minta ellenőrzése, kizárólag részleges garanciát tud adni arról, hogy a szervezet kontrolljai a vezetőség által elvártan operálnak (BARTA 2018e).

Az audit lebonyolításának egyszerűsített folyamatát szemlélteti 6 lépésben az alábbi, 2. ábra DAVIS et al. (2011) alapján.



2. ábra: Az auditálás folyamata

Forrás: DAVIS et al. (2011)

- **Tervezés:** Az audit tervezési fázisában az auditorok az adott szervezet vezetésével közösen meghatározzák az audit hatókörét, mely történhet rotációs formában, eseti alapon, vagy kockázatelemzés alapján a legkockázatosabbnak ítélt terület vizsgálatával. Az audit hatóköre tartalmazhat bizonyos informatikai rendszereket, hálózatokat, telephelyeket és irodákat, üzleti folyamatokat és kiterjedhet csak a kontrollkörnyezet egy bizonyos területére is pl. hozzáféréskezelés.

- **Helyszíni vizsgálat:** A tervezés végeztével és annak jóváhagyását követően kezdődik a helyszíni vizsgálat, melyben az auditorok különböző audit technikákat alkalmazva (pl. folyamatok megfigyelése és reperformálása, interjúk szakterületi vezetőkkel, stb.) megbizonyosodnak a tesztelt kontrollok tervezeti, implementációs és működési hatékonyságáról.
- **Megállapítás és validáció:** A helyszíni vizsgálat során, amennyiben az auditorok hiányosságot vélnek felfedezni a vizsgált hatókörben, azt egyeztetik a szakterületi vezetőkkel, felméri a hiányosság kockázatát és egy auditjelentés formájában dokumentálják észrevételeiket.
- **Javaslattevés:** A javaslattevés fontos részét képezi az auditori munkának, mivel az auditor szükségszerűen a felfedezett kontrollhiányosságokra magas szinten alkalmas javaslatokat megfogalmazni, amivel az eltérések kijavíthatók, ezzel fejlődik a szervezet biztonsági környezete, a javaslatok megfogadásával az audit elnyerte célját.
- **Riportolás:** A riportolási fázisban az audit csapat bemutatót tart a vezetőség részére a feltárt hiányosságokról és hivatalos, szakterületi válaszokkal ellátott riport-formában a vezetés és befektetők részére bocsátja azt.
- **Nyomon követés:** Az auditori munka folyamatos, az auditort szükséges bevonni a riportolás után is, ahol az auditor feladata az észlelt eltérések nyomon követése, azaz, hogy a hiányosságokat kijavítani célzó korrekciós akciótervek megvalósultak-e határidőre és sikeresen képes volt a szervezet a fennálló kockázatokat mérsékelni.

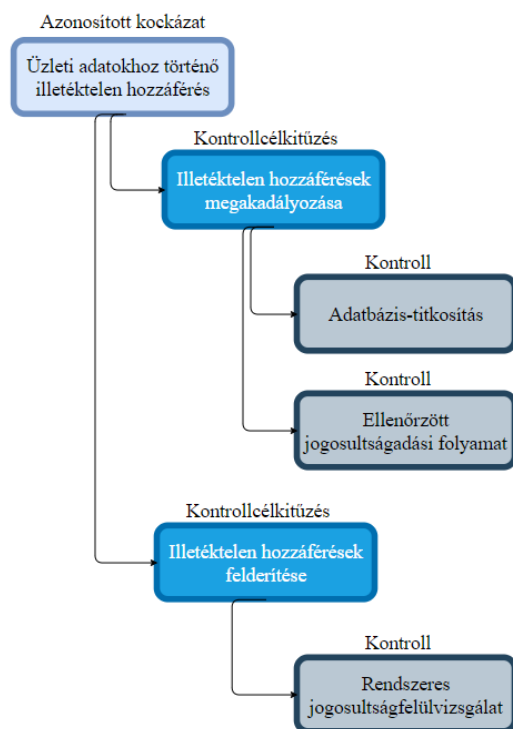
Az ISACA (Information System Audit & Control Association 2016) észlelési kockázatnak nevezi az informatikai auditban azt a kitétséget, amikor a kontrolltesztelési módszertanok nem megfelelőek, azaz nem képesek egy adott informatikai kockázatot mitigáló kontrollok hatékonyságának felmérésére, ezáltal torzképet adva annak helyes működéséről. Gyakorlati példával illusztrálva, annak ellenőrzése, hogy egy vizsgált rendszer biztonsági konfigurációi az audit ideje alatt ideális, tehát az auditor egy időpontra meghatározva vizsgálja a beállításokat, magában hordozza azt a kockázatot, hogy az auditált fél az audit idejére megváltoztatta azt a követelményeknek megfelelően, majd az audit végeztével egy gyengébb konfigurációt állít be ismételt, mivel az üzemeltetési szempontból egyszerűbbnek tekinthető pl. technikai felhasználók esetén a jelszavas védelem hiánya. Az audit tervezésének időszakában, ennek megfelelően, az auditor feladata az észlelési kockázat minimalizálása is átgondolt tesztelési tervek definiálásával.

Mindazonáltal kiemelendő és a dolgozat szempontjából is kritikus, hogy az auditor nem-megfelelő szakmai ismerete, csupán látszólagos függetlensége, esetleg hanyagsága is az észlelési kockázat növelésével jár, tehát az automatizált objektív auditálás kényszer az emberi hibákból eredő kockázatok csökkentésére, mely megalkotása a dolgozat legfőbb eredményterméke (4.2. alfejezet).

2.2.4. Kontrollok tesztelése

Az audit célkitűzése a kontrollok ellenőrzése annak megbizonyosodására, hogy az implementált kontrollok alkalmasak a szervezet által detektált kockázatok mérséklésére. A kockázatok

csökkentése érdekében szükséges a kontrollcélkitűzések definiálása, egy adott kontroll akkor tekinthető működőképesnek, ha eléri a kitűzött célokat (3. ábra).



3. ábra: Kockázatok, kontrollcélkitűzések és kontrollok kapcsolata

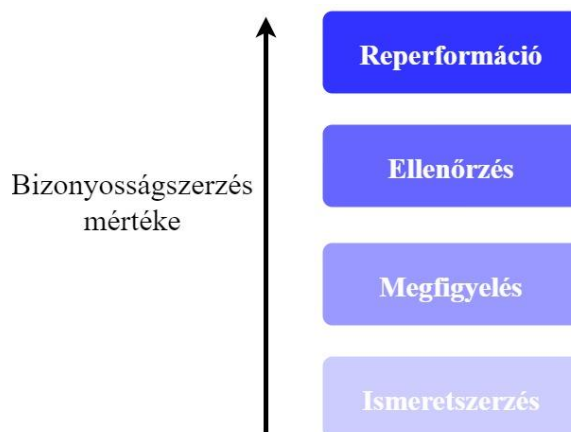
Forrás: Saját szerkesztés

Kontrollok tesztelése alapvetően három elkülönült szinten valósulhat meg (Information System Audit & Control Association 2015):

- **Tervezet/kialakítás:** Az auditor megvizsgálja a kontroll terveket, szabályzatokat és leírásokat, hogy azok tervezett megvalósításai alkalmasak-e kockázatok mitigálására. Amennyiben az auditor egy folyamatleírás alapján azt tapasztalja, hogy pl. az auditált szervezet tervezetten a mentési lemezeket a szerverszobában tárolja, úgy a kontroll sérül, mivel egy esetleges tűz esetén nem csak a szerverek veszhetnek kárba, hanem a mentések is (tehát a redundáns rendelkezésre állás nem biztosított), ezért a folyamatot újra kell gondolni, mégpedig, a mentéseket egy külső helyszínre kell szállítani rendszeres időközönként.
- **Implementáció:** Az auditor, feltéve, hogy adott kontroll tervezeti szinten megfelelő, értékeli annak bevezetését, azaz a szervezet a gyakorlatban implementálta-e a kontrollt, vagy az csak szabályzati szinten „működik”.
- **Működési hatékonyság:** Az audit ellenőrzi, hogy az implementált kontroll folyamatosan, egy időszakra levetítve is hatékonyan operált, nem történt kivétel. Abban az esetben, ha a mentések heti szinten elszállításra kerülnek, de egy bizonyos időszakban erről a felelős vezető megfigyelkedett, akkor, bár a kontroll implementált, működési hatékonysága megkérdőjelezhető.

A fenti példákból következtethető, hogy a három vizsgálati szint egymásra épül. Amennyiben egy kontroll tervezet szinten nem megfelelő, az auditor feleslegesen tekinti annak megimplementációját, nem alkalmazható kockázatot csökkentő óvintézkedésként, továbbá, az implementáció sérülésével a működési hatékonyság nem mérhető.

Az auditorok a megbizonyosodás mértékének elvárásai szerint 4 különböző szinten végezhetnek vizsgálatokat (4. ábra).



4. ábra: Audit technikák a bizonyosságszerzés mértékének rangsora szerint

Forrás: BARTA – GÖRC SI (2021)

- **Ismeretszerzés:** Az auditor a kijelölt vezetőkkel párbeszédet kezdeményez, szakmai interjút szervez, vagy egyéb írásos dokumentum vizsgálatával győződik meg a kontroll jelenlétéről. Észlelhetően ez a megbizonyosodás mértékének legalsóbb szintje, mivel a rendelkezésre álló dokumentumok és munkatársak, akaratlanul vagy akaratlagosan, az igazsággal nem megegyező információt állíthatnak.
- **Megfigyelés:** Az auditor helyszíni vizsgálatot tart és saját érzékszerveivel bizonyosodik meg a kontrollok megfelelőségéről pl. szerverszoba viziten vesz részt.
- **Ellenőrzés:** Az auditor a kontrollok működésének alátámasztásául megerősítő evidenciát (auditbizonyítékot) kér be az auditált féltől pl. könyvvizsgálat esetén bizonylatokat ellenőriz.
- **Reperformáció:** A megbizonyosodás legfelső foka. Az auditor az auditált fél által végzett folyamatot megismétli, és ezáltal nyer tanúbizonyságot annak helyességéről pl. újra könyvelési számlák egy adott csoportját a mérlegeredmény alátámasztására.

Az észlelési kockázatot, tehát a kontrollfolyamatok megfelelőségének vizsgálatára irányuló eljárások a megbizonyosodás mértékétől függően befolyásolhatják.

2.2.5. Az ISO:IEC 27001:2013 szabvány

A belső információbiztonsági környezet felépítését nem kell szükségszerűen a vezetésnek és alkalmazott szakemberiknek teljesen zérusról kiviteleznie, számos keretrendszer és ajánlás áll rendelkezésre, mely segítséget nyújthat a kontrollok kialakításában, úgymint a korábban említett

CobIT 2019 (HAES et al. 2018), NIST 800-53 (2013) vagy ISO (2013) szabványok. Az egyik legszélesebb körben alkalmazott standard, melyet akkreditációs szervezet is tanúsíthat az ISO/IEC 27001. 2018 év végén az ISO (International Organization for Standardization - Nemzetközi Szabványügyi Szervezet) által végzett kutatás alapján, világszerte 31910 vállalat rendelkezett bejegyzett tanúsítvánnyal, míg ez a szám magyarországi viszonylatban 484 volt (ISO 2018).

A szabvány két jól elkülöníthető részre osztható (ISO/IEC 27001 2013):

- **Információbiztonsági Irányítási Rendszer (4.-10. fejezet követelményei):** Az információbiztonsági szervezet és a biztonság, mint irányítási rendszer folyamatainak együttese. Többek között magában foglalja a belső szabályozói környezet meglétét, fenntartását és ellenőrzését, a kockázatmenedzsment folyamatok kialakítását, az auditálás szükségességét, valamint a folyamatos fejlesztési tevékenységek összességét;
- **Kontrollcélkitűzések és kontrollok (A melléklet):** A szabvány által előírt kontrollterületenként osztályozott kontrollkövetelmények és a kontrollcélrendszer együttese. A melléklet 14 kontrollterületet és 114 kontrollt nevez meg a 2. táblázatban leírt rendszerben.

A szabvány erőssége, hogy rugalmasan szervezetre szabható, azaz a szervezet szempontjából irrelevánsnak tekinthető kontrollok figyelmen kívül hagyhatók, továbbá, addicionális kontroll megléte esetén az a szabvány ésszerű felépítésének és jól strukturáltságának köszönhetően az beépíthető.

2. táblázat: Az ISO/IEC 27001:2013 szabvány kontrollterületei

Hivatkozási szám	Kontrollterületek megnevezése	Kontrollkövetelmények száma
A5	Az információbiztonság vezetői irányítása	2
A6	Az információbiztonság szervezete	7
A7	Humán-erőforrás biztonsága	6
A8	Vagyon-menedzsment	10
A9	Hozzáférés szabályozás	14
A10	Titkosítás	2
A11	Fizikai és környezeti biztonság	15
A12	A működtetés biztonsága	14
A13	A kommunikáció biztonsága	7
A14	Rendszer beszerzés, fejlesztés és karbantartás	13
A15	Szállítói kapcsolatok	5
A16	Információbiztonsági incidensek kezelése	7
A17	A működésfolytonosság információbiztonsági aspektusai	4
A18	Megfelelőség	8
Összesen		114

Forrás: ISO/IEC 27001:2013 A melléklet

A kontrollterületek rövid bemutatása a 2. számú mellékletben található az ISO/IEC 27001:2013 alapján.

2.2.6. Az alfejezet összefoglalása

A saját kutatási téma vonatkozásában, az információról, információbiztonságról és auditról folytatott szakirodalomkutatás az alábbi pontokban erősítette meg kutatói munkám létjogosultságát:

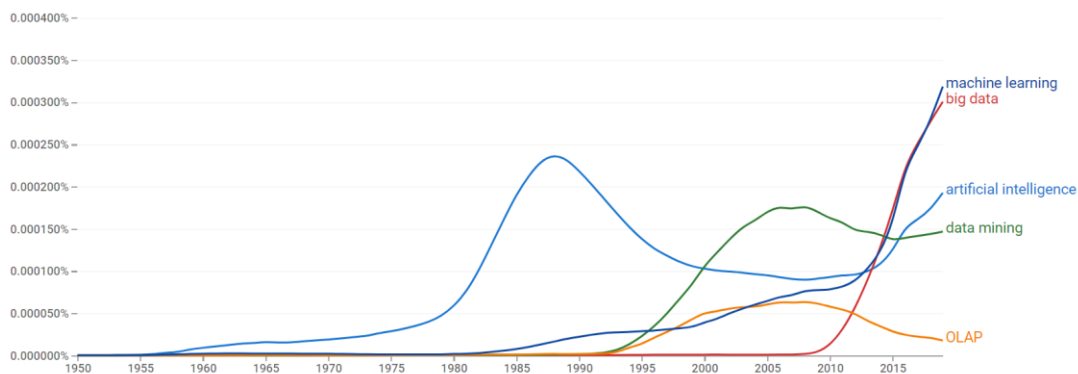
- Az információ, az információrendszerekben tárolt adatok kritikus értékkel bírnak a szervezetek számára, így azokat védeni kell a külső behatolókkal, jogosulatlan munkatársakkal és a természeti katasztrófákkal szemben;
- A szervezeti információbiztonsági szint kontrollokkal fokozható, azonban a kontrollok működési hatékonyságát folyamatosan nyomon kell követni. Az implementálandó kontrollokat és azok szintjét az információbiztonságra irányuló informatikai kockázatelemzés, mint üzleti támogató folyamat képes meghatározni, mely szakmai terület képviselői az audit osztály munkatársai;
- A kontrollok nem kizárólag elemi szinten, hanem együttesen is funkcionálhatnak, ezért a kontrollkörnyezetet több dimenzió mentén érdemes vizsgálni. A kontrollok közötti kapcsolatok feltárása és összefüggések ellenőrzése komplex munkát igényel, ami a dolgozatban a mesterséges intelligencia kapcsán kerül újszerűen kezelésre.

Mivel az audit manuális eszközökkel képtelen abszolút bizonyosságot szerezni a kontrollkörnyezet hibamentes működéséről, mely részben köszönhető annak, hogy statisztikai mintaelemszámokkal történik az ellenőrzés, valamint mindenkor fennáll az észlelési kockázat valószínűsége, szükségesnek ítéltető egy olyan automatizált megoldás (mesterséges intelligencia), mely képes kockázati alapon becslést adni egy kontrollhiányosság meglétére, mely az audit minőségi munkájának visszaellenőrzésére irányul.

2.3. Mesterséges intelligenciával ellátott modellfejlesztés

2.3.1. A mesterséges intelligencia fogalmi megközelítései

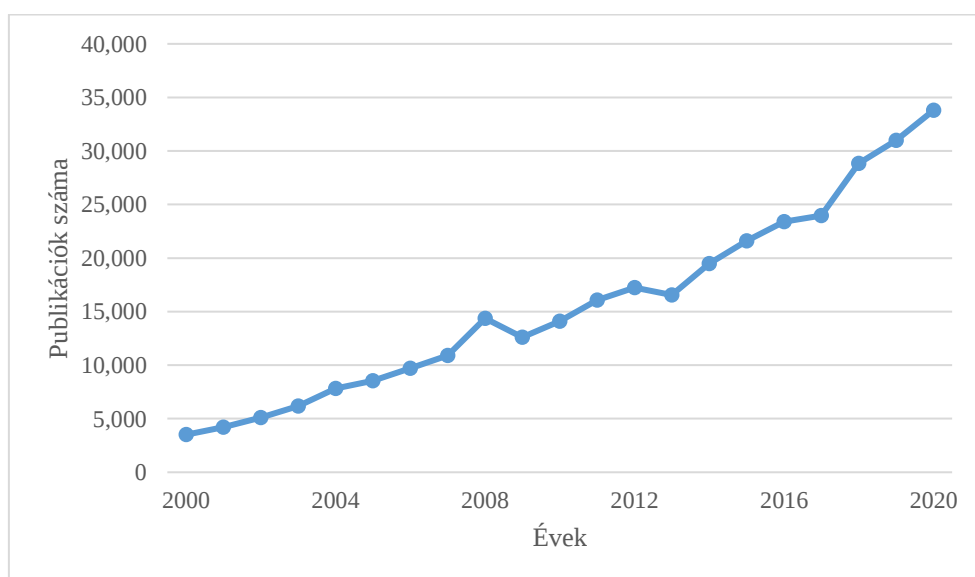
A mesterséges intelligencia és ahhoz szorosan köthető fogalmak és a tudományterületet átfogó definíciók, algoritmusok és alkalmazási területek nem számítanak új keletűnek a kutatók számára. Az 5. ábra szemlélteti a Google adatbázisában tárolt dokumentumokban megtalálható, az ábrán leolvasható kifejezések népszerűségét (Y tengely) az előző 70 évben (X tengely).



5. ábra: A mesterséges intelligenciához szorosan kapcsolódó fogalmak népszerűsége

Forrás: Google Books Ngram Viewer (2020)

A mesterséges intelligencia területének kutatása hullámzó, jelenleg fokozott népszerűségnek örvend, mely a kutatási témában megnövekedett publikációk számából is igazolható (6. ábra), többek között köszönhető a számítógépek grafikus kártyáinak megnövekedett kapacitásának és feldolgozó képességének (LIU et al. 2017), s nem mellesleg a kvázi kiszámíthatatlanul változó, jelenleg masszívan adott/növekvő társadalmi befogadóképességnek (vö. MI-koalíció, Ipar 4.0).



6. ábra: „Mesterséges Intelligencia” kulcsszót tartalmazó publikációk száma a Scopus-ban

Forrás: Scopus (2021a)

A fogalom, mint „*mesterséges intelligencia*” megalkotása John McCarthynak köszönhető, aki 1956 nyarán két hónapos munkatalálkozót szervezett a számításelmélet és atomelmélet amerikai kutatóinak számára Dartmouthban. A munkatalálkozó előkészítő anyagában kiemelte a mesterséges intelligencia önálló tudományterületté válásának szükségességét (MCCARTHY et al. 1955).

Ugyanakkor, a mesterséges intelligenciához köthető első igazán jelentős elmélet és gyakorlati alkalmazás megszületése, véleményem szerint, WARREN STURGIS MCCULLOCH és WALTER PITTS (1943) nevéhez fűződik, akik 1943-ban publikálták a perceptron modellt (MCP-neuron). Az agysejtek funkcionalitását alapul véve, egy egyszerűsített bináris klasszifikációs eljárást ajánlottak, mely az agy biológiai működésének analógiájára hivatott gépiesíteni a tanulási folyamatot. Ez volt a legelső modell, mely a mai nap is ismert neurális háló alapjait képezi, és innen indult el annak széleskörű kutatása. Kiemelendő, hogy ezen az elven működik a jelenleg legmodernebbnek feltételezett technológia a „deep learning” (magyar szakirodalomi fordításban: mély neurális háló), melynek alkalmazási területei kiterjednek az önvezető autók, virtuális asszisztensek, virtuális valóságok, automatikus játékgépek és ajánló rendszerek fejlesztésére is (BARTA 2018a, BARTA – GÖRCSI 2019). DIAMANT (2016) álláspontja, hogy jelen korunk mesterséges intelligencia kutatásának fejlődése nagyrészt a neurális hálóknak köszönhető. NG (2018) a mesterséges intelligencia forradalmát köti a „deep learning” fogalomkörhöz, míg VARGA és CSEH (2019) a mesterséges intelligenciát a robotizáció forradalmának nevezi.

A mesterséges intelligenciát az évek során számos kutató definiálta. RUSSEL és NORVIG (2005, 2009) ezen definíciókat rendszerezte, és arra jutottak, hogy a mesterséges intelligencia definícióit két dimenzió mentén értelmezik a terület kutatói. A megfogalmazások csoportosítását a 3. táblázat szemlélteti.

3. táblázat: A mesterséges intelligencia definícióinak csoportosítása

	Emberi teljesítmény	Racionalitás
<i>Gondolkodó rendszerek</i>	Emberi módon gondolkodó rendszerek	Racionálisan gondolkodó rendszerek
<i>Cselekvő rendszerek</i>	Emberi módon cselekvő rendszerek	Racionálisan cselekvő rendszerek

Forrás: RUSSEL - NORVIG (2005, 2009)

A táblázat felső része a gondolkodásra, az észszerűsége helyezi a hangsúlyt, míg az alsó része a viselkedésre, a konkrét cselekvésre. A táblázat bal oszlopa az emberi képességet és teljesítményt helyezi előtérbe, a jobb oldala a racionalitást. Az emberi teljesítményt és racionalitást érdemes két részre bontani, mégpedig azért, mert az ember nem feltétlen mindig racionális döntést hoz, ami a számítógépek esetén követelmény.

A két dimenzió mentén, a mesterséges intelligencia fogalmait, így négy különböző részre lehet osztani:

- Emberi módon gondolkodó rendszerek
- Racionálisan gondolkodó rendszerek
- Emberi módon cselekvő rendszerek

- Racionálisan cselekvő rendszerek

A mesterséges intelligencia efféle optimalizációt nélkülöző és kizáró kategorizálása, véleményem szerint, nem felel meg a tudományos kutatásokra irányuló mesterséges intelligenciával ellátott modellezési követelményeknek, a Knuth-i elvnek, valamint merőben megkérdőjelezhető a kategorizálás egyértelmű elkülöníthetősége pl. létezik-e cselekvés gondolkodás nélkül? A szakirodalomban jelen lévő definíciók értékelése, azonban azt mutatja, hogy a fogalomalkotások elhelyezhetők az említett dimenziókban. Megjegyzendő, amennyiben szükséges a dolgozatban ismertetett pozitivista megközelítést és modellezési gyakorlatokat elhelyezni a mátrixban, úgy az a racionálisan gondolkodó és cselekvő rendszerek kategóriába tartozna a leginkább.

Az emberi módon gondolkodó rendszerek fogalomkörhöz tartozik többek között, BELLMAN (1978), HAUGELAND (1985), BORGULYA (1998), SHABBIR és ANWER (2018) definíciói. A mesterséges intelligencia:

- „az emberi gondolkodással asszociálható olyan aktivitások automatizálása, mint pl. döntéshozás, problémamegoldás és tanulás” (BELLMAN 1978)
- „újszerű kísérlet, hogy a számítógépeket gondolkodásra készítsük” (HAUGELAND 1985)
- „olyan módszerek, amelyek az emberi problémamegoldást, következtetési folyamatot, vagy heurisztikus megközelítéseket modelleznek” (BORGULYA 1998)
- „ma rendelkezik az emberi intelligencia képességeivel, utánozza azt, különböző feladatokat végez, amelyekhez gondolkodás és tanulás szükséges” (SHABBIR - ANWER 2018)

A racionálisan gondolkodó rendszerek fogalomkörhöz tartozik többek között CHARNIAK és MCDERMOTT (1985), WINSTON (1992) és ARAÚJO (2014) definíciói. A mesterséges intelligencia:

- „a mentális képességek tanulmányozása számítási modellek segítségével” (CHARNIAK - MCDERMOTT 1985)
- „az észlelést, a következtetést és a cselekvést biztosító számítási mechanizmusok tanulmányozása” (WINSTON 1992)
- „a klasszikus statisztikai-matematika önkorlátozó jellegének feloldása révén olyan új ismeretszerzési és ismeretábrázolási formák kidolgozása, melyek segítségével régóta ismert feladatok új megvilágításba kerülhetnek, s eddig megoldatlan problémák megoldhatóvá válnak” (BUNKÓCZI 1998)
- „képes az alapvető logikai következtetések levonására” (ARAÚJO 2014)

Az emberi módon cselekvő rendszerek fogalomkörhöz tartozik KURZWEIL (1990), RICH et al. (2009), MATA et al. (2018) definíciói. A mesterséges intelligencia:

- „olyan funkciókat teljesítő gépi rendszerek létrehozásának a művészete, amelyhez az intelligencia szükséges, ha azt emberek teszik” (KURZWEIL 1990)
- „annak tanulmányozása, hogy hogyan lehet a számítógéppel olyan dolgokat művelni, amiben pillanatnyilag az emberek a jobbak” (RICH et al. 2009)
- „egy olyan kiterjesztett tudományterület, mely lehetővé teszi a számítógépek részére, hogy problémákat oldjanak meg komplex biológiai folyamatok emulálásával, mint a tanulás, érvelés és önkorrekció” (MATA et al. 2018)

A racionálisan cselekvő rendszerek fogalomkörhöz tartozik többek között POOLE et al. (1998), NILLSON (1998), DILEK et al. (2015), YU és KUMBIER (2017), PFEFFER et al. (2017) definíciói. A mesterséges intelligencia:

- „az intelligens ágensek tervezésének a tanulmányozása” (POOLE et al. 1998)
- „a műtárgyak intelligens viselkedésével foglalkozik” (NILSSON 1998)
- „lehetővé teszi számunkra, hogy olyan autonóm számítástechnikai megoldásokat tervezzünk, melyek alkalmazkodni tudnak felhasználási környezetükhöz” (DILEK et al. 2015)
- „lényegében adatvezérelt. Az ember-gép együttműködésén keresztül statisztikai koncepciókat fogalmaz meg, és így adatokat generál, algoritmusokat fejleszt és értékeli az eredményeket” (YU - KUMBIER 2017)
- „képes automatizálni feladatokat” (PFEFFER et al. 2017)

A „modernebb” definíciók jelentős része összeköttetésbe hozza a mesterséges intelligenciát az automatizálással és az adatvezérelt döntéshozatallal, míg a korábbi, főleg 2012 előtti, megfogalmazások futurisztikusabb hangvételt teremtenek. Ennek oka, hogy a korábban is említett Big Data feldolgozása mesterséges intelligencia eszköztárával együtt magas színvonalú tudás kiépítésére képes, tehát a technológiák párhuzamos fejlődése egymás húzó erejével is rendelkeznek. A minőségi Big Data megléte követelmény (és ahhoz kapcsolódó tudásreprezentációs technikák), mely az alapvető tudást szolgáltatja a mesterséges intelligencia matematikai feldolgozó egységeinek.

A mesterséges intelligencia szakirodalmi definícióiban ellentétek is meghúzódnak. Míg PFEFFER et al. (2017) az automatizálásra helyezi a hangsúlyt, elvonatkoztatva az emberi intelligencia reperformálásának képességétől a mesterséges intelligenciát, addig egyes kutatók, többek között MATA et al. (2018), SHABBIR és ANWER (2018) teljes egészében az emberi intelligenciával asszociálják azt, mely szerintem helytelen. PFEFFER et al. (2017) gondolatait megalapozó, DIAMANT (2016) érdekes módon azt részletezi, hogy a mesterséges intelligencia kutatásoknak nem feltétlenül kell az emberi biológia és pl. az agy működését szimulálnia, nem az az elsődleges vizsgálandó terület, mivel az amőbák és baktériumok is képesek intelligens cselekvéseket produkálni, úgy, hogy biológiai aggyal nem rendelkeznek.

Álláspontom szerint, mivel számos olyan algoritmus is létezik, melyek képesek kvázi intelligenciát szimulálni, és az intelligencia alatt itt és most megalapozott racionális döntést értek, az emberi/természetes intelligencia és a mesterséges intelligencia párhuzamosításának nincs feltétlen létjogosultsága.

Hangsúlyos kiemelni a mesterséges intelligencia teremtőképességét, melynek alapjait GOODFELLOW et al. (2014) fektette le a GAN (Generative Adversarial Network) neurális hálón alapuló innovatív megközelítésével, mely képes az eredeti adathalmazzal statisztikailag egyenértékűt előállítani, tehát egy alternatív valóságot létrehozni. A GAN lehetővé teszi a hamis valóság megalkotását, mely életre keltette a „DeepFake” szintetikus közeget, mely számos információbiztonsági sérülékenységhöz vezethet (pl. visszaélés személyes adatokkal, hamis hírek automatikus generálása, stb.). A GAN képes könyvek és zenei kompozíciók írásának automatizálására, mely HADJERES et al. (2017) kísérlete alapján olyan jól teljesít, hogy még a szakértőket is meg lehet téveszteni, hogy melyik az emberi és melyik a gép által kreált mű.

Az emberi cselekedetek sajátosságaival felruházó mesterséges intelligencia definíciók, véleményem szerint, az emberi hasznosság céleszközeként gondolnak a mesterséges intelligenciára. Álláspontom, hogy a mesterséges intelligenciát el kell vonatkoztatni az emberi/természetes intelligenciától, mivel az ember elfogult, szubjektív, és intuitív módon hoz döntéseket. Az a véleményem, hogy nem az emberi intelligencia lemodellezése, azaz gépiesítése a cél, hanem egy racionalitást támogató objektív, az emberi belemagyarázó képességtől mentes, pártatlan eszközt kell, hogy kínáljon, azaz az emberiség munkáját és intelligenciáját kiegészítő, erőforrás-optimalis megoldásokkal kell, hogy szolgáljon. Ennek okán a mesterséges

intelligenciának szükségszerűen magában kell foglalnia az emberi tudás legjavát, így nem lehet teljesen emberi kötődés nélküli, azonban a gépekben rejlő intelligencia mibenlétét a felsorakoztatott érvek szerint szükséges értelmezni.

KNUTH (1995) alapján „Tudomány az, amit értünk annyira, hogy elmagyarázzuk egy számítógépnek. Minden más művészet”. Knuth híres mondatát PITLIK et al. (2017) az alábbiak szerint értelmezi: „Tudás/tudomány az, ami forráskódba átírható – minden más emberi aktivitás művészet, ahol a tudomány és a művészet egymáshoz képest nem értékelhető – lévén ezek ugyanazon érme két oldalát jelentik.” Az, hogy valami már Knuth-i vagy sem, még csak a vízvázlasztó, mert a nem Knuth-i világ mellett a már Knuth-i világon belül létezik a mesterséges intelligencia megoldások versenye, vagyis a „jó<jobb<legjobb” keresése és a mesterséges intelligencia a céltalanság tételének kreatív feloldási kísérleteivel jellemezhető leginkább (PITLIK 2014). Véleményem szerint az idézetek a tudományos kutatás objektivitására hívják fel a figyelmet, azaz az emberi szubjektivitás kiküszöbölésére, tehát olyan modellek megalkotása a mindenkorai cél, mely a legtöbb emberi részrehajlástól mentesített bemenetet kapja. PITLIK et al. (2017) ezt kontextus független (context free) jellegű modellezésnek nevezi, mely hermeneutikája kiterjed a teljes kombinatorikai térre. PITLIK et al. (2017, 2020a, 2020c) munkáiban gyakorta megjelenik a konzisztencia fogalma, melyet úgy definiál: „A hasonlóságelemzések önellenőrzésének egyik rétege, mely azonos probléma/modellkérdés esetén egymással logikai kapcsolatba hozható párhuzamos modellek összevetése kapcsán elvárt rendezettségére utal. Konzisztenciaalakzat tetszőlegesen sokféle lehet. Teljesen konzisztens (ellentmondásmentes) eredmény formálisan nincs, de lehet olyan kevés konzisztencia-modellréteg, melyek eredményei nem ütnek egymást.” (PITLIK 2014). A fogalom, értelmezésem szerint kiemeli, hogy a modellezésben mindenkor elvárt az önerősítő mechanizmusok visszaellenőrzése, azaz akkor járhat a modellalkotás sikerrel, ha annak részegységei logikailag is következetes döntésre jutnak, mely felfogás az alapjait képezi a dolgozatban ismertetett kutatásnak és megfeleltethető a Knuth-i elvnek és pozitivisták szemléletmódnak.

Az ismertetett definíciók és gondolatok értékelése alapján az alábbi kritériumot rendelem a mesterséges intelligencia fogalomalkotásához:

- vonatkoztasson el az emberi szubjektivitástól (működjön minél inkább kontextus független módon);
- törekedjen minél komplexebb célfüggvény esetén az optimális megoldás megtalálására a rendelkezésre álló erőforrások függvényében;
- az emberi üzemszerű észlelőképesség/hermeneutika meghaladásával, legyen képes a rejtett mintázatok felismerésére;
- maximalizálja a kinyerhető információt az elérhető adatokból;
- a kinyert információ értékességi aspektusai és hibái hassanak vissza a következő hasonló problémamegoldási folyamatra.

Feltételezhetően, olyan gép és matematikai apparátus nem áll kiforrott formában nyilvánosan rendelkezésre a mai nap, amely az összes kritériumnak maximálisan eleget tenne (vö. céltalanság tétele). Azért is limitált a mesterséges intelligencia mai vonatkozásában, mert az adat jelenléte kényszer a tökéletes modellalkotáshoz, azonban az adat csak korlátokkal képes leírni a világot – de ezek a korlátok is egyre inkább felismerhetők. Mindazonáltal, a dolgozat célja, olyan modellek megalkotása, melyek képesek a legközelebb kerülni az általam felállított, az elfogadott szakirodalom által alátámasztott követelményekhez.

A mesterséges intelligencia algoritmusokra általánosságban elmondható, hogy számításigényes folyamatok összessége, tehát, bár létezett korábban tudományos elmélet egy adott szakterületi

probléma megoldására, azt a gyakorlat kevésbé tudta követni és kevésbé tudta igazolni az elméletek létjogosultságát a hiányzó technológiai háttér miatt. Neurális hálók, ahogy fentebb említésre került, már a 40-es években is léteztek koncepció szintjén, de jelen korunk fejlettsége adja a lehetőséget, hogy az azokban rejlő számítási erőt kihasználjuk, és olyan alkalmazásokat hozzunk létre, melyek sokáig csak elméletben léteztek. Erre a legjobb példa az önvezető autót támogató szoftverkomponensek, melyek mély neurális háló alapon működnek (MAQUEDA et al. 2018).

Az értelmezett és kielemezett szakirodalmi definíciók alapján, a következő saját megfogalmazást alkalmazom, melynek mindenkori követelményét a fentebb felsorolt kritériumok tükrében szükséges értelmezni:

A mesterséges intelligencia a racionális döntések sorozatát hivatott támogatni, olyan matematikai eszköztárt kínálva, mellyel a rendelkezésre álló adatvagyonból lehetséges a logikus következtetések levonása és számítási problémák optimalizálása, kontextus független módon a konzisztens modellezés követelményeit felhasználva.

S ezen definíció értelmében az összeadástól a neurális hálókig a mesterséges intelligencia maga az adatfeldolgozás.

2.3.2. Gépi tanulás

A gépi tanulás, véleményem szerint, a mesterséges intelligencia fogalomalkotásának egy potenciális teljesítményjavító kutatási területe, ezért kényszerűen magában foglalja azt. A gépi tanulás célja a rendelkezésre álló adathalmazban meghúzódó mintázatok automatizált felismerése, azok matematikai leképezése (BARTA 2018c). A gépi tanulás RASCHKA (2015) megfogalmazása alapján: „*a számítógépek felruházása a tanulás képességével*”. Sokkal technikaibb definíciót ad meg RUSSEL és NORVIG (2009): „*A tanulás alap gondolata az, hogy a megfigyeléseket ne csak az ágens jelenlegi cselekvéseinek kialakítására használjuk, hanem arra is, hogy javítsuk a cselekvésre való jövőbeli képességeit*”. A definícióban megjelenik a jövő, azaz a predikció megnevezése, miszerint a gépi tanulás egy olyan folyamat, mely hozzájárul az ismeretlenre vonatkozó sejtések előrejelzéséhez. DUA és DU (2011) szerint a gépi tanulás „*egy tudományos modell építése, mely képes a jelenlegi adatokból tudást képezni*”, tehát a tudás létrehozásával asszociálja a gépi tanulást. HASTIE et al. (2009) egyszerűen „*az adatokból való tanulás*” mikéntjeként definiálja a fogalmat. NG (2018) szimplán úgy fogalmaz, hogy a gépi tanulás „*A rendelkezésre álló adathalmaz attribútumaihoz történő célváltozó rendelése egy függvényen keresztül*”. CHOLLET (2018) új szoftverfejlesztési paradigmának nevezi a gépi tanulást (7. ábra).

A klasszikus programozási megközelítés szabályalapon operál, azaz előzetesen definiált feltételekhez köti a cél elérésének útvonalát, ahol adott feltételrendszerhez történő megfelelés vagy nem-megfelelés esetén a program különböző kimenetekhez irányítja a felhasználót. Belátható, hogy ez a módszer hosszútávon nem fenntartható. Egyrészt, komplex problémák megoldása nem kivitelezhető szabályok megalkotásával, például, egy arcfelismerő rendszerhez manuálisan a pixelek alapján véges sok szabály megírása embert próbáló és időigényes, nem beszélve arról, hogy a fényviszonyok, a kamera dőlésszöge, a személy aznapi megjelenése és egyéb külső tényezők az összes szabály megalkotását ellehetetlenítik. Másrészt, a programozó nem feltétlen ismeri az összes szabályt, vagy az emberiség nincs adott tudás birtokában, mely hozzájárulna a

rendszer jóságához. Erre világít rá DUA és DU (2011) megfogalmazása, tehát a gépi tanulás egyik kitüntetett célja a tudás létrehozása, mely az adatok közötti ok-okozati összefüggések feltárása, ami nem feltétlen ismert előzetesen. Véleményem szerint, ezért sincs létjogosultsága az információt és tudást egymásnak megfeleltetni, mint ahogy az kifejtésre került a 2.2.1. alfejezetben.



7. ábra: A klasszikus programozás és gépi tanulás paradigmája

Forrás: CHOLLET (2018)

CHOLLET (2018) a szabályrendszer „megteremtését” a gépre bízta, legyen az képes a saját mechanizmusai révén az alkotásra, úgy, hogy az adatokat és ismert válaszokat bemenetként a gép rendelkezésére bocsátjuk. Összefoglalva, a gépi tanuló eljárásoktól elvárjuk az explicit szabályalkotás-nélküli operációt.

A gépi tanulás során rendelkezésre álló válasz lehet egy fénykép esetén a személy megnevezése, mely által a gép képes lehet a személyt a későbbiekben azonosítani egy másik fényképen. Azonban, a CHOLLET (2018) által definiált paradigma nem tér ki azokra az esetekre, amikor nincs egyértelmű válasz pl. hasonlóságelemzés, klaszterezési eljárások, stb. ahol ugyancsak az adatban rejlő tudás felderítése a cél matematikai leképezésekkel, így a CHOLLET (2018) szerinti gépi tanulás inkább a felügyelt tanulásnak feleltethető meg, mely a következő alfejezetben (2.3.2) kerül bemutatásra.

Számos példa mutatkozik arra, hogy a tudás kinyerése nem minden esetben lehetséges a gépi tanuló algoritmus által alkalmazott mechanizmustól függően. HOLZINGER et al. (2019) fekete doboz algoritmusoknak nevezi azokat az eljárásokat, amelyek bemenete és kimenete ismert (értelemszerűen), azonban a módszer belső működése transzparencia hiányában nem engedi meg a „miérték” feltárását, az adatfeldolgozó szerkezet által végzett műveletek és/vagy adatttranszformáció nem enged betekintést a következtetési logikába. A fehér doboz modellek, ellenkezőleg, magyarázatot adhatnak a kimenetek létjogosultságára. Fekete doboz algoritmus, többek között a neurális háló és SVM (support vector machines – szupport vektor gépek), valamint az együttes módszerek döntő többsége, míg fehér doboznak minősülnek, általánosságban, bár típustól eltérően, a döntés fák.

Bizonyos problémák esetén a következtetési logika megértése kritikus. Gondoljunk csak a gyógyszergyártásra, ahol a háttér biológiai és élettani folyamatok megértése nélkülözhetetlen egy adott gyógyszer biztonsági minősítéséhez (pl. BIRÓ – PITLIK 2020). Egy másik szemléletes példa a fizikai jelenségek magyarázata. Ha a világtörténelem elismert fizikusai évszázadokkal ezelőtt fekete doboz modellezéssel vizsgálták volna az Univerzumot, akkor lehetséges, hogy kifogástalanul modellezni tudták volna az égitestek keringési pályáit, de nem érthetnénk, hogy azok miért ellipszis alakú pályákon keringenek, valamint a fizika általános törvényeiről is számottevően kevesebb tudás állna rendelkezésünkre. Hitelebírálásnál sem elégednénk meg azzal a magyarázattal, hogy az ügyfelek által leírt attribútumok közötti mintázatok miatt nem vagyunk

jogosultak hitelt felvenni, kíváncsiak lennénk, hogy milyen tulajdonságot szükséges „javítani” ahhoz, hogy a jövőben pozitív legyen a bírálat eredménye, és gyakran az sem egyértelmű, hogy a rendelkezésre álló attribútumok milyen irányba mutatnak a jóhoz. Bár a „mit kellene javítani?” kérdésre a látszatkorrelációkat/mintázatok értelmező szimulátorok továbbra is képesek választ adni – valódi ok-okozatiság nélkül – s ez a pragmatizmus visszaköszön a genetikai algoritmusokban és majd a dolgozat keretei között a genetikai potenciál által vezérelt keresési stratégiákban is (4.3. alfejezet).

Láthatóan számos gyakorlati alkalmazhatósági példa szól a fekete doboz modellezés ellen, azonban az vitathatatlan, hogy megannyi területen (pl. kép-, beszéd és hangfelismerés) a fekete doboz algoritmusok kiemelten jobban teljesítenek, mert nem vállalnak fel olyan restriktiókat, melyek mellett nem lehet jobb a modellezés, mint amilyen. Egy arcfelismerő alkalmazás esetén könnyedén eldönthető, hogy adott megoldás működőképes-e (felismer egy adott személyt) és a felhasználó által érdektelen lehet, hogy melyek voltak a képen azok a minták, mely alapján beazonosíthatóvá vált.

A gépi tanuló modellezésben felhasznált független változók (attribútumok) a célváltozóra tett pozitív hatását az adott attribútum preferált irányának lehet nevezni, más szóval iránypreferenciának. A hitelbírálásnál egyértelmű: minél nagyobb a hitelfeltevő havi keresete (ceteris paribus), annál hitelképesebb változatlan hitelösszeget feltételezve. Azonban az iránypreferenciák meghatározása korántsem triviális, valamint nem feltétlenül állnak lineáris/monoton kapcsolatban a célváltozóval, sőt, az attribútumok kölcsönös hatásainak is lehet együttes iránypreferenciája, melynek optimuma létezhet, így a modellezésekben polinomhatással/periodicitással is kell számolni. Az iránypreferenciák meghatározása, letapogatása és visszaellenőrzése (vagyis a ceteris paribus alakzatok karakterisztikájának hitelessége a fehér doboz felé való elmozdulásként) kritikus a gépi tanuló modellezésben, mivel az ideális modell megalkotásához tisztában kell lenni azzal, hogy egységnyi attribútum adott irányba történő elmozdítása milyen hatással bír (ha egyáltalán bír) a kívánt célváltozó értékeire. A dolgozatban a modellezési technikák bemutatásánál és kiértékelésénél nagy hangsúlyt fektettem az iránypreferenciák értelmezésére és érvényességvizsgálatára.

Felismerve az iránypreferenciákból adódó kockázatokat, az is kijelenthető, hogy egy adott modell attribútumainak érték-irány megismerése a modelljósághoz vezető út egyik meghatározó ismérve. A gépi tanulás kimeneteként adott rendszerválaszok hasonlóan rendelkezhetnek érték-iránnyal, amennyiben a rendszerválaszok tulajdonságait több dimenzió mentén definiáljuk. Ez azt jelenti, hogy a modellek jósága a rendszerválaszok minősítése által becsülhetővé válik, így vélelmezhető, hogy klasszikus performancia metrikák előzetes ismerete nélkül is van lehetőség a jó megismerésére.

2.3.3. Gépi tanuló rendszerek típusai

A gépi tanuló eljárások kategorizálását a szakirodalom többé-kevésbé egységesen kezeli. HASTIEE al. (2009) a gépi tanulást felügyelt és felügyelet nélküli (nem felügyelt/felügyeletlen) tanulásként csoportosítja, míg RASCHKA (2015) egy harmadik kategóriát is külön kiemel, ami a megerősítéses tanulás (8. ábra).



8. ábra: A gépi tanulás típusai

Forrás: RASCHKA (2015)

- A **felügyelt tanulás** egy meglévő adathalmaz, azaz historikusan összegyűjtött információ, alapján történő mintázat-keresési eljárások összessége, mely a mintázat feltárását követően hozzájárul egy adott célváltozó értékének előrejelzéséhez (SAGAR 2015). Felügyelt tanulási eljárók közé sorolhatók az osztályozók (klasszifikálók), valamint a regressziós módszerek, ill. a dolgozatban is alkalmazott hasonlóságelemzés termelési függvény-generáló potenciálja. Az osztályozási problémák esetén a célváltozó nominális (kategória) változó, míg a regressziós eljárások egy metrikus változót becsülnek. Osztályozási feladatnak tekinthető egy kézírásfelismerő rendszer, ahol az egyes diszkrét kimeneti értékek megfeleltethetők az ábécé karaktereinek. Regressziós probléma lehet pl. egy adott település lakásárainak prediktálása.
- A **felügyelet nélküli** tanuló algoritmusok alkalmasak az adathalmaz struktúrájának és szerkezetének feltárására, ahol nincs dedikált célváltozó (vagy az látens/fiktív). A mintázat leképezésének célja a hasonló tulajdonsággal rendelkező adatrekordok csoportosítása, valamint a normától történő eltérés vizsgálata. A klaszterező algoritmusok és a hasonlóságelemzés anti-diszkriminatív potenciálja tipikusan felügyelet nélküli eljárások. A marketingkutatásban gyakorta használatosak a piaci szegmentációk feltárására, így azok elemzése és kiértékelése a felügyelet nélküli tanulási eljárások egyik gyakori alkalmazása. Felügyelet nélküli tanulásnak nevezhetők, továbbá, a dimenziócsökkentő eljárások pl. Főkomponens analízis (SZELÉNYI 2001). Ezekben az esetekben az eljárások/módszerek értékét a matematika mibenléte (pl. az optimalizálás) adja, mely tovább erősíthető a filozófiai/hermeneutikai rétegekkel.
- **Megerősítéses tanulásban** a rendszer feladata, hogy az (ágens) képes legyen a környezetéhez alkalmazkodni és optimális stratégiát válasszon egy adott veszteségfüggvény minimalizálása, vagy jutalomfüggvény/célfüggvény maximalizálása érdekében (PETER - NORVIG 2009). Az automatizált sakkjáték a megerősítéses tanuláson alapuló gépi tanuló rendszer egy klasszikus példája – de ide sorolható a (pl. hasonlóságelemzési, illetve hibrid) modell-láncok konzisztencia-maximalizáló jelenségek is.

A gépi tanuló rendszerek, egy másik aspektusból vizsgálva, lehetnek egyszerű/naiv/alap/gyenge technikák és együttes alkalmazások (ensemble methods). Az együttes alkalmazások egyfajta meta-tanulók, melyek az alapló módszerek közös kombinációja, együttes mechanikája alapján törekednek a jószágmetrikák javítására (HEARTY 2016). Számos architektúra elképzelhető az alaptanulók hibrid felhasználására, melyek lehetnek pl. szekvenciálisan egymásba ágyazva, párhuzamosan alkalmazva aggregált döntéshozásra, stb. A teljesség igénye nélkül, az alábbi megközelítések tekinthetők népszerűnek, mint együttes módszerek - s ezek léte/megszületése átvezet a klasszikus statisztikai matematika szemléletmódjából a mesterséges intelligencia-alapúság világába, hiszen

az alaptanulók olyan elemek a matematikai statisztika világából, melyek kombinálása új logikákat/módszertanokat vár el:

- **Boosting (fokozás/turbózás):** A technika az alaptanulók szekvenciális összekapcsolását teszi lehetővé, ahol a soron követő algoritmus az előző hibájából képes tanulni pl. az adathalmaz újraszűrésével. A Boosting módszerek alapvetően a tanulás hatékonyságának növelését szolgálják a tanulómintában jelenlévő „anomáliák” pl. kiugró értékek súlyozásával, melyek következtében a variancia (eltérés a tanulómintán és tesztmintán mért teljesítménymutatók között) csökkentése nem elsődleges cél (GÉRON 2017). Boosting technikára példa az AdaBoost és Gradiens Boosting (Gradient Boosting), melyeket a dolgozatban is felhasználok az üzleti probléma megoldására (4.2. alfejezet), mivel a kutatás elsődlegesen magas pontosságot kíván elérni.
- **Bagging (Bootstrap Aggregation – Bootstrap Aggregálás):** A módszer a rendelkezésre álló adathalmaz attribútumait és objektumait (általában véletlen) mintavételezéssel szelektálja, majd alaptanulókat illeszt a kiválasztott adathalmazokra. Az egyes alapmodellek becsléseit aggregálja, mely lehet pl. többségi szavazás (módusz), mely így a variancia csökkentését célozza a tanulási függvény minél pontosabb leképezése helyett. Bagging módszerre példa a Random Forest (Véletlen Erdő) algoritmus (BARTA 2018c).

Az együttes módszerek alkalmasak stabil, gyakran magas általánosító képességgel rendelkező modellek létrehozására és gyakorlatilag bármely alaptanuló esetén felhasználhatók, mindazonáltal, számításigényes eljárások. A szakirodalomban tetten érhető, hogy az együttes módszerek többségében magasabb performancia elérésére képesek összehasonlítva azokat az alaptanulókkal pl. DELGADO et al. (2014), BARTA (2018c), BARTA – PITLIK (2020). A magasabb performancia ára általában a nagyobb komplexitás, a hosszabb futásidő, mely a valós idejű vezérlő/szabályozó rendszerek esetén korlátként értelmezendő, de a gyártási (alapanyag és kapacitás) költségek esetén is el lehet jutni a gazdasági értelemben már „nem éri meg” kategóriák határát.

A gépi tanuló rendszerek a becslés approximációjára alkalmazott eljárás karakterisztikáit vizsgálva, lehetnek parametrikus vagy nem-parametrikus technikák (DANGETI 2017). A parametrikus módszerek diszkriminatív függvény előállítását/becslését kísérelik meg, ahol a tanulás eredménye a függvény együtthatóinak letapogatása. Erre példa a neurális háló, logisztikus regresszió, Naive Bayes, stb. A nem-parametrikus modellek, ezzel ellentétben, más eljárás/metrika mentén képesek az összefüggések leképezésére. Pl. a KNN (K-nearest neighbour – K-legközelebbi szomszéd) az adatpontok közötti hasonlósági alapon egy meghatározott távolság-metrikát alkalmazva hoz döntést arról, hogy a nem ismert rekord a leghasonlóbb objektumokat alapul véve, mely pl. osztályba sorolható. Nem-parametrikus módszer még az SVM, döntési fák, stb. Mint látható, a hasonlóság egy olyan centrális absztrakció, melyre a matematikai statisztika és a mesterséges intelligencia is visszanyúl – lévén ez a jelenség maga az általánosító képesség motorja.

2.3.4. Gépi tanuló rendszerek fejlesztése

A gépi tanuló rendszerek fejlesztése az információrendszerek fejlesztéséhez hasonlóan több fázisra/csoportra bontható, mely fázisok/csoportok tervezése azért indokolt, hogy a lehetséges döntési pontok (pl. adattranszformáció, modellválasztás, metrikaválasztás, stb.) a modell objektíven mért jóságához a lehető legnagyobb mértékben hozzájáruljon. A paramétercsoportok

azonban végső soron együtt fejtik ki hatásukat – a leíró jellegű csoportosításra, mint minden rendszerező tudáskezelésre csak az emberi megértés potenciáljának növelése érdekében van szükség.

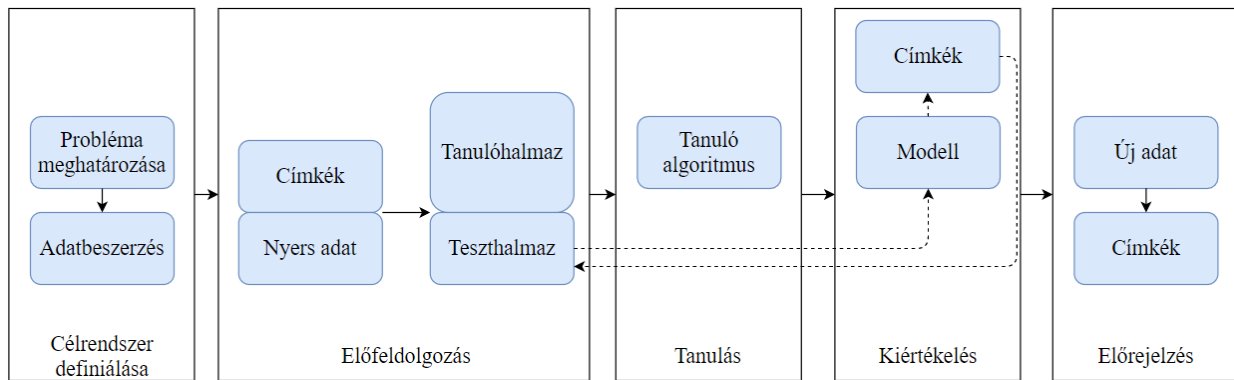
A gépi tanuló eljárások fejlesztése esetén célszerű az általános alkalmazás- és információrendszer fejlesztési módszerekkel magas szinten párhuzamot vonni, azonban álláspontom szerint, ezt ki kell terjeszteni a tudományos modellek tervezésének és fejlesztésének módszertanával. Gépi tanuló rendszerek esetén is igaz az állítás, hogy az alkalmazást bevált iparági gyakorlatok és kutatási eredmények által igazolt módszertanok mentén érdemes fejleszteni. SZEPESNÉ (2010) az információrendszerek fejlesztését 3 részre osztja.

- Első szakaszban, melyet SZEPESNÉ (2010) előszakasznak nevez, történik a kiindulási helyzet elemzése, a feladat megfogalmazása, a költségelemzés, illetve az előnyök definiálása, tehát az információrendszer célkitűzésének meghatározása, valamint az iránypreferenciák értelmezése, a jó (célfüggvény) fogalmának megalapozása;
- A második szakasz a fejlesztés szakasza, melyben kivitelezésre kerül az adatbázis struktúra leírása, az alrendszerekre való bontás, a specifikációk véglegesítése és a rendszerkomponensek együttműködésének biztosítása;
- A harmadik szakaszban, a felhasználói szakaszban, veszi kezdetét az alkalmazás használata, karbantartása, felülvizsgálata és optimalizálása.

KOVÁCS (2009) ajánlást fogalmaz meg a tudományos szimuláció és modellépítéssel kapcsolatban.

- Elsőként a modellalkotás célját (pl. a jó fogalmát) kell meghatározni;
- Második lépés a vizsgálandó rendszer definiálása;
- Harmadik lépés a modellkalibráció, azaz az inputként szolgáló adathalmazt szükséges a modellre szabni;
- Negyedik lépés a validáció (pl. a jóságot leíró komplex pl. konzisztencia-alapú skála alkalmazása), annak ellenőrzése, hogy a modell elérte-e az előzetesen definiált célját.

A felhasznált szakirodalmak áttekintésének segítségével az alábbiakban leírtak szerint érdemes megtervezni és fejleszteni a kutatási célok elérését szolgáló gépi tanuló alkalmazásokat, melyben figyelmet kell szentelni az iparági gyakorlatok és a tudományos céllal történő modellalkotás elvárásainak. Az általam alkalmazni kívánt modellalkotási folyamat a következő, melyet vizuálisan a 9. ábra szemléltet RASCHKA és MIRJALILI alapján (2019), mely leginkább a felügyelt gépi tanuló rendszerek karakterisztikáit írja le, valamint az ábra kiegészítésre került KOVÁCS (2009) és SZEPESNÉ (2010) gondolataival.



9. ábra: Gépi tanuló rendszerek fejlesztése a prediktív modellezés keretében

Forrás: RASCHKA – MIRJALILI (2019), KOVÁCS (2009) és SZEPESNÉ (2010) alapján

1. A fejlesztés első fázisa a gépi tanuló rendszer célkitűzésének definiálása. Ebben a szakaszban történik az adott üzleti probléma kiértékelése, és annak meghatározása, hogy az új alkalmazás milyen módon lesz képes az üzleti problémát megoldani, elhárítani, karbantartani. A célkitűzés meghatározásakor a felhasználni kívánt adatvagyon, annak beszerzési kritériumait, a függő és független változókat is meg kell határozni. Amennyiben a specifikáció megfelelő minőségű eredményeket produkál, úgy a következő pont az adatok beszerzése. Adatokat lehetséges különböző forrásokból összegyűjteni, mely lehet saját terepmunkán gyűjtött adatok rögzítése, adatok vásárlása harmadik személytől, az adatok munkahelyen belüli kollektívja, internetes letöltése, stb. Az adatok begyűjtése után az adatok transzformációjára van szükség, azaz az előfeldolgozásra.
2. Mivel az alkalmazni kívánt adathalmaz gyakorta nem áll megfelelő formában rendelkezésre, így azok tisztítása, transzformációja, tehát előfeldolgozása indokolt. Azon felül, hogy egy adott gépi tanuló modell teljesítményét is képes fokozni, a magas minőségű adattranzformáció ellenállóbb lehet külső támadások ellen, ahol az alkalmazás funkcionalitásának megkárosítása a támadó célja (BHAGOJI et al. 2018). Az adatok előfeldolgozása esetén az alábbi kihívásokkal szembesülhet a kutató:

- **Hiányzó értékek az adathalmazban:** HASTIE et al. (2009) felveti, hogy hiányzó értékek esetén egyik lehetséges megoldás a hiányzó függő változó oszlopában szereplő adatok átlagát, móduszát vagy mediánját venni, majd a hiányzó értékeket azzal helyettesíteni. RASCHKA (2015) szerint, sok esetben a hiányzó rekord teljes törlése is indokolt lehet, mivel a hiányzó adatokkal való adatfeldolgozás a modell predikciós erejét csökkentheti, összefüggéseket indukálhat. BEALUC és ROSENTHAL (2018) egy klasszifikációs és regressziós döntési fából álló algoritmust javasol, mely kutatási eredményeik alapján magasabb teljesítménnyel szolgál, mint az adatok átlagolása, vagy egyéb statisztikai módszerekkel való kitöltése. Modelljükben a hiányzó adatok feltöltését egy másodrendű gépi tanuló eljárásnak kezelik, ahol a hiányzó értékeket kell előrejelezni. A hiányzó értékek esetén mérlegelendő, hogy az adatok eltávolítása mekkora kárt tehet a teljes adathalmaz értelmezhetőségében, amennyiben az csak egy kitüntetett függő változóra jellemző, érdemes lehet kizárólag a változó elhanyagolása. A kutatásban alkalmazott adatvagyon javarészt anonim volt, így számos értékes információ nem állt rendelkezésre, melynek részleteit a 3.4. alfejezet tárgyalja. Az adathiányos pozíciók teljes értelmezési intervallummal való lefedése (kombinatorikai

alapon) egyben a közgazdaságtanban is ismert érzékenységvizsgálat jelenségének modellezésre való kiterjesztését jelenti, ahol pl. az irány-preferencia-stabilitásának zavarai is felismerhetők a Hartman-elvet meghaladó módon (PITLIK et al. 2020b).

- **Különböző mérési skálán mért adatok transzformálása:** Az adathalmaz a céltól függően különböző mérési skálán szereplő adatokat tartalmazhat, melyek lehetnek nominális (névleges), ordinális (rendezéses), intervallum (különbség), vagy arányskálán mért adatok (SZŰCS et al. 2008). A nominális és ordinális skálán mért adatok szövegszerű értékeket tartalmazhatnak, ezért ezen adatok kódolása szükséges, hogy a számítógép képes legyen az adatfeldolgozásra (BODON 2010). CERDA et al. (2018) a „one-hot-encoding” módszert ajánlja, melynek lényege, hogy a szöveges/nominális értéket tartalmazó függő változókból új változók létrehozásával, a változó különböző értékeinek számával megegyezően, orvosolni lehet a problémát, ahol a változó csak ott vesz fel pl. 1 értéket, ahol az az eredeti változóban adott érték jellemző volt rá. A nominális skálák információ-értékének optimalizált kezelésére nyújt megoldást pl. a hasonlóságelemzés exploratív modell-rétege, ahol a formálisan létező lépcsősfüggvények flexibilisen polinomizálódnak a nominális skálaértékek variációnak megfelelően.
- **Adattranzformáció a modell teljesítményének növelése érdekében:** HACKELING (2014) felhívja a figyelmet, hogy különböző gépi tanuló módszerek érzékenyek, ha az adatok nem ugyanazon a skálán szerepelnek. Ilyen pl. a neurális hálók, logisztikus regresszió, SVM, KNN algoritmusok, melyek performanciája nagyban függ az alkalmazott adattranzformációs eljárástól (MCCLURE 2017). Az adatok transzformációja lehetséges, többek között az adatok standardizálásával vagy normalizálásával, melynek egy speciális nézete a rangsorszámozás, mely látszólag információvesztéssel jár és a hasonlóságok és a lépcsős függvények kapcsán ez a kockázat előnybe fordul át.
- **Függő változók statisztikai transzformálása:** A függő változók nagyszáma performancia problémához vezethet, azon túl a változók értelmezése is megnehezedhet, így SAJTOS és MITEV (2007) a sok homogén jellemzővel rendelkező adathalmazban a változók közötti összefüggések feltárására a faktorelemzést javasolja, mely a változók redukcióját eredményezi. A kutatásban felhasznált adatvagyon esetén nem éltem ezen lehetőséggel (faktorelemzés), azonban a hasonlóságelemzések esetén a mindenkor függő változó transzformációja szinte kötelező, mert a futás-optimalizálást eltolásokkal, nagyításokkal lehet csak garantálni annak érdekében, hogy a módszertan a pozitív egész számok körében maradjon, ami a lépcsős függvények számára ideális.
- **Függő változók kiválasztása:** Az adatgyűjtés során olyan adatok kerülhetnek be az adathalmazba, melyek egyáltalán nem járulnak hozzá a célváltozó előrejelzéséhez, ezért azok kvázi használhatatlanok az üzleti cél elérése érdekében. Zavaró jelek a független változók között is lehetnek, de ezek felismerése, kizárása vagy éppen integrálhatóságának kikényszeríthetősége mind olyan modellezési rugalmasság, amire szükség van (ZHENG – CASARI 2018). Másrészt, a párhuzamos függőváltozókra készülő párhuzamos modellek a konzisztencia-elemzések alapjai. A függő változó lehet maga az idő és a tér is (PITLIK et al. 2005).

- **Adathalmaz felbontása tréning-, teszt-, és validációshalmazra:** Ahhoz, hogy megfelelően validálni lehessen az elkészítendő modell performanciáját az adatokat tréning-, teszt-, és validációshalmazra érdemes bontani. A tréning adatok szolgálnak a modell tanítására, vagyis a mintázatok felfedezésére és azok elraktározására. A tesztalmazon, mely független a tréninghalmaztól, elvégzett számítások adnak becslést a modell általánosító képességére, mivel olyan adatokat tartalmaz, melyet a modell korábban nem látott. A validációshalmaz alkalmazható a kiválasztott modellek további finomhangolására, vagyis a belső paraméterek optimális kombinációjának megtalálására. Az adathalmaz efféle megbontása egy megkerülhetetlen hátránnyal jár: a tanuló algoritmust értékes adatvagyonról fosztjuk meg, mely így nem épül be a tanulásba.
 - **Irány-preferenciák meghatározása:** Különösen a rangsorszámmal operáló algoritmusok esetén az irány-preferenciák előzetes meghatározása, érvényesség vizsgálata és visszaellenőrzése szükséges. A tetszőleges irány-preferenciák alkalmazása mellett a konzisztencia fogalma speciális, kontextus független értelmet nyer. Ennek extrém esete a függvény-szimmetriák sérülésének felhasználása validációs célra, valamint a magas „hamis pozitív” találati aránnyal rendelkező modellek korlátozására.
3. A gépi tanuló eljárások adatfeldolgozó szakaszát követően a fejlesztési (tanulás) szakasz veszi kezdetét. A szakaszban kiválasztásra kerülnek az alkalmazni kívánt modellek - kivéve, ha a modellalkotás maga már olyan flexibilis, hogy tetszőlegesen tud váltani az ismert modell-típusok és ezek hibridizálása között lényegében a változók és matematikai művelet, szintaktikai jelek kombinatorikai terében keresve a nem letális modell-variánsokat minél nagyobb hatékonyság mellett. A dolgozatban alkalmazott modellek elemzésében és kiválasztásában a kontrollhiányosságok detektálása, mint elsődleges cél van fókuszban. A kontrollhiányosságok detektálása egy gyanúgenerálási probléma (lásd 2.4. alfejezet), mely többféle szempontból megközelíthető. Értelmezésem szerint, felfogható, mint egy klasszifikációs/osztályozási feladat, melyben a gyanú egy külön osztályként jelenik meg, tehát a modellezésekben olyan gépi tanuló eljárásokat szükséges alkalmazni az output becslésére, melyek bizonyítottan alkalmasak klasszifikációs feladatok elvégzésére. Másodsorban, hasonlóság alapon, a gyanúmomentumok csoportosításának is van értelme (pl. azok elkülönítése a normális magatartásoktól), tehát empirikus kutatással vizsgálandó feladat a különböző csoportok/klaszterek felderítése. Harmadsorban, a nem ismert gyanúmomentumokat kiugró értéként is lehet kezelni, melynek következtében statisztikai módszerekkel is van lehetőség a gyanúgenerálásra, azonban ezt a megközelítést a dolgozat mellőzni kívánja. Mindezek technikai megvalósítását szem előtt tartva, a modellépítés során alkalmazott kimeneti struktúra eltérő lehet, így olyan algoritmusokra van elsődlegesen szükség, melyek a kívánt struktúrák leképezésére alkalmasak. Ugyanakkor, az ismert technikák hibridizálására is van lehetőség, lehetséges a különböző algoritmusokat a hibrid modell részfeladatainak optimalizálására is alkalmazni az egyes döntési pontok és feladatok megoldásában, melynek célja az egyes gépi tanuló módszerek előnyeinek és erősségeinek kiaknázása, melyek más algoritmusok esetén gyengeségnek bizonyulnak. A dolgozatban élek a hibrid modellezés adta lehetőségekkel, mivel az erősségek kihasználása potenciális jószágmetrika javulással járhat.

4. A modellalkotást követően azok kiértékelését kell elvégezni annak érdekében, hogy a teljesítményt mutatószámokkal ki lehessen fejezni. A jószág kritériumokra előzetesen metrikákat kell építeni, hogy objektivitása biztosítható legyen. A dolgozatban alkalmazott jószágmetrikákat a 3.3.1. alfejezet összegzi. A jószág fogalma nem lehet egysíkú a céltalanság tétele értelmében, mert az egyes önálló jószág-rétegek között létezik függetlenségi tér (PITLIK 2014).
5. A modell építése, valamint a jószágmetrikák megfelelő teljesítményének elfogadása után a gépi tanuló rendszer éleskörnyezetbe való állítása következik (előrejelzési fázis), melyben a rendszer üzemszerű működése veszi kezdetét, ahol a valós idejűség jellege a mindenkori hardver-erőterek alapján lényegi elvárás. Megjegyzendő, hogy a dolgozat sajnos nem tudott kilépni számos szervezési lépés ellenére sem superszámítógépes infrastruktúrára.

2.3.5. A modellalkotás kihívásai és kockázatai

A gépi tanulás lehetővé teszi az explicit szabályalkotás-nélküli következtetési logikák megteremtését, mely a tradicionális szoftverfejlesztési paradigmát kiváltva/kiegészítve, hozzájárul a komplex, jól körbe határolt szabályrendszert kevésbé alkalmazható problémák megoldásához (vö. indukció).

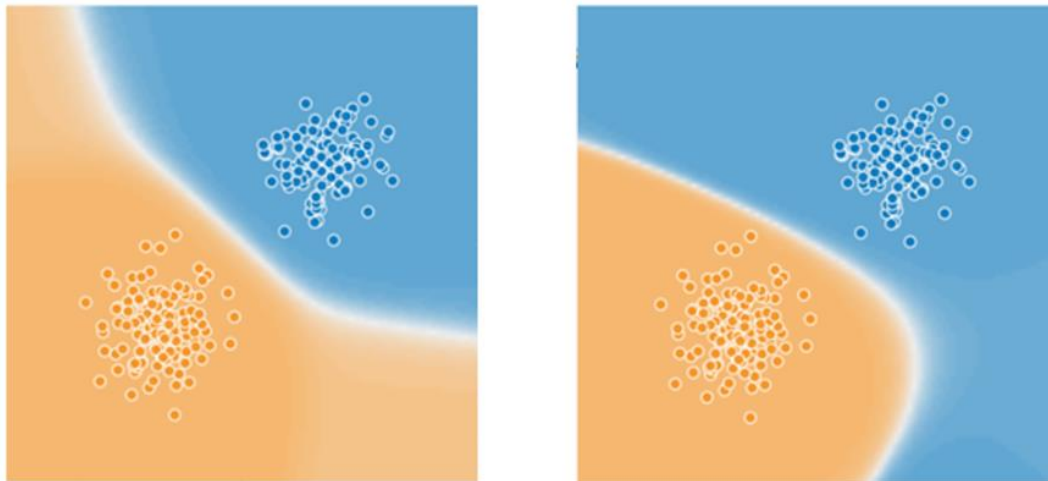
A gépi tanulásra képes rendszerek, azonban, számos esetben kudarcot vallanak, ami a rendszer fejlesztésének és üzemeltetésének egyes fázisaiban bekövetkező kockázati faktorokra vezethető vissza. Ezen tényezők kihívások elé állítják a kutatókat és szervezeteket egyaránt. Korábbi kutatásaim alapján, összefoglalóan a teljesség igénye nélkül, az alábbi kockázatok azonosíthatók a gépi tanuló projektekben annak életciklusa szerint rendezve (BARTA – GÖRCSI 2018, 2020, 2021):

- Célrendszer definiálás:
 - Adatok rendelkezésre állásának hiánya és korlátozott beszerezhetősége, valamint költségvonzata;
- Előfeldolgozás:
 - Reprezentativitást nélkülöző attribútumok bevonása a tanulási folyamatba;
 - Elégtelen/túlzott adattisztítás;
 - Adatminőségi diszkrepancia (az adatminőség optimalizált ellenőrzésének hiánya);
 - Tanuló, validációs és tesztadathalmaz nem megfelelő (mennyiségi és minőségi) felosztása, a tanulásba bevont adathalmaz teljességi hiánya;
 - A jószág fogalmának zavarai;
- Tanulás:
 - Célszerűtlenül alkalmazott modellezési eljárások;
 - Programozási és logikai hibák;
 - Algoritmusok optimalizálásának hiánya;
- Kiértékelés:
 - Irreleváns jószágmetrikák és performancia mutatók alkalmazása;
 - Alul- és túlilleszkedés, modell-egyensúlyok elemzésének elmaradása;
 - Problémaspecifikus kiértékelés (pl. kiegyensúlyozatlanság kérdésköre);
 - Kimenetek helytelen/nem-automatizált/ráérzés-alapú értelmezése;

- Előrejelzés:
 - Fejlesztési-, teszt- és az éles üzemkörnyezet közötti technológiai eltérések, kompatibilitási problémák;
 - Általános biztonsági kockázatok (pl. jogosultsági kérdések);
 - Rendszer nyomon követésének és az adatbázis frissítésének hiánya;
- Egyéb:
 - Adatok felhasználásának jogi akadályai;
 - Az alkalmazás etikai kérdéseket vet fel;
 - Nem áll rendelkezésre szakértő a fejlesztéshez/üzemeltetéshez;
 - Benchmark-adatok hiánya.

Korábbi kutatásaim és kísérleteim arra is rávilágítanak, hogy a gépi tanuló rendszerek esetén a rendelkezésre álló adatok elérhetősége és minősége jelenti a legkritikusabb pontját a gépi tanulás megvalósításának és annak alkalmazhatóságának (BARTA – GÖRCSI 2018, 2020, 2021). Ennek miéért könnyen belátható, mivel az adat szolgáltatja az információt, végső soron a tudást, ezáltal képes a bizonytalanság csökkentésére. Kiemelendő, nem a több adat a jobb, hanem a több minőségi adat, így az adatgyűjtés követelménye a dolgozatban a minőségi adatbeszerzésre korlátozódik. Belátható, hogy irreleváns és/vagy a célváltozó értékét torzító tanulóhalmaz káros, így indokolt lehet a tanulóhalmaz redukálása a jószág ideálisabb approximációjához.

Az alábbi ábrán (10. ábra) egy osztályozási probléma látható két különböző döntési határral (BARTA 2017). A két osztály egy neurális hálóval lett különválasztva különböző aktivációs függvényeket felhasználva (az első tangens, a második logisztikus). A két osztály láthatóan, tökéletesen szeparálható lenne lineárisan is, azonban az elválasztás helye korántsem triviális, mivel az adatok által leírt tér nem teljes, tehát egy újabb adatpont megjelenése (pl. a bal felső sarokban) az egyik, akár mind a kettő modellt érvénytelenítheti, pontosságát csökkentheti.



10. ábra: Osztályozási probléma különböző döntési határokkal

Forrás: BARTA (2017)

Az adatok beszerzésére, gyűjtésére és feldolgozására számos megközelítés alkalmazható, ezek karakterisztikái az alábbiak:

- **Több minőségi adat gyűjtése:** Az adathalmaz kibővítése hozzájárulhat a magasabb teljesítményű modellezéshez, azonban ez számos esetben kivitelezhetetlen, függően az alkalmazás hatókörében lévő szakterületi problémától. Egy korábbi munkámban (BARTA et al. 2020) a GDPR alapján elemeztem a büntetések sajátosságait, azonban, amelyik országban nem volt büntetés, nem lehetett levonni konklúziót. Ilyen esetekben a tanuló rendszernek az elérhető adatokból kell dolgoznia, a beszerzés lehet, hogy kivitelezhetetlen, komplex vagy túlságosan drága (pl. beruházás igényes). A publikusan elérhető adatok is korlátosak, valamint főbb jellemző tulajdonságuk, hogy oktatás jelleggel elérhetőek, ezért valós probléma megoldására gyakran alkalmatlanok.
- **Adat-augmentáció:** Adat-augmentációs technikákkal növelhető az adatmennyiség és csökkenthető a kényszerű túlilleszkedés. Az adat-augmentáció központi eleme a rendelkezésre álló adatok minimális változtatása pl. egy kép elforgatása vagy árnyékolása. Általában kép- és hangfelismerő és elemző rendszerek tanításánál alkalmazott módszer, ezért más alkalmazási területeken, ahol emberi érzékszervekkel (látás, hallás) nem triviális a kimenet, nem alkalmazható (SHORTEN – KHOSHGOFTAAR 2019).
- **Szintetikus adatok:** A szintetikus adatok olyan mesterségesen előállított adatok, melyek statisztikailag valós adathalmazra hasonlítanak pl. a korábban is említett GAN alkalmas szintetikus adatok létrehozására. A szintetikus adat, azonban, nem a valós világot írja le, így hasonlóan az adat-augmentációhoz, főleg képfelismerő alkalmazásoknál használatos módszer (SHORTEN – KHOSHGOFTAAR 2019). Mindamellet megjegyzendő, hogy az adatvagyon mesterséges objektumainak léte alapjaiban képes meghatározni a feltáruló modellek tulajdonságait. Ilyen speciális objektum lesz a disszertációban a mindenkori genetikai potenciált legjobban közelítő mesterséges/kereséssel feltárt objektumok sora.
- **Transzfertanulás:** A transzfertanulás lehetőséget ad egy modell „tudását” pl. parametrikus modellezés keretében az együtthatókat, hasonló céllal fejlesztett modellek között transzferálni, azaz ismételten felhasználni, melynek kiemelt felhasználási területe szintén a képfelismerés alkalmazási területei (YANG et al. 2020). Legnagyobb kihívása, hogy a megoldandó problémáknak hasonló céllal és kialakítással kell bírniuk (pl. önvezető autók személygépjármű felismerő modulja, mely felhasználható egy sebességmérő alkalmazásban), valamint, a transzferált tudás, mivel egy adott modell-architektúrából érkezik, így a célrendszeren is hasonló architektúrával kell, hogy rendelkezzen (pl. neurális hálók esetén ugyanannyi rejtett réteggel és neuronnal).
- **Mikromodellezés:** A mikromodellezés egy adott probléma részegységeinek felbontását jelenti, külön modellt építve a jól elkülöníthető részfeladatokra. Például a járművek felismerésére irányuló rendszer esetén külön algoritmus osztályozhatja azokat személygépjárművekre, teherautókra, stb.

Egy, a kutatás témájában szereplő alkalmazás fejlesztése esetén, a fentebb felsorolt eljárások nem alkalmazhatók:

- **Több adat gyűjtése:** Az adatvagyon bizalmassága végett, nagyon korlátos a beszerezhetősége.

- **Adat-augmentáció:** Az adatvagyon módosítása nem ad valós képet a megismerni vélt világról.
- **Szintetikus adatok létrehozása:** Hamis adatok létrehozása nem ad valós képet a megismerni vélt világról.
- **Transzfertanulás:** A legjobb tudomásom szerint nem elérhető publikus adat a kutatás témájában, melynek kényszerűen hasonló struktúrában kellene rendelkezésre állnia.
- **Mikromodellezés:** A probléma, tekintve az adatok korlátozott jellegét, nem bontható szét jól strukturálható részfeladatokra.

Véleményem szerint, a probléma karakterisztikáit vizsgálva, olyan megoldásokra van szükség, melyek képesek az elérhető adatok optimális-közeli felhasználására a gépi tanulás genetikai potenciálját kiaknázva. Továbbá, a tesztelésre felhasznált adatok csökkentik a tanulás sikerességét, mivel értékes információ vonódik el, amit a modell beépíthetne a tanulási mechanizmusba, ezért létjogosultsága van olyan technikák kutatásának, melyek képesek a maximális tudást kinyerni az adatvagyonból, valamint úgy megtalálni a modellek között a legideálisabbat, hogy halmaz-szeperáció ne legyen szükséges. Ezen lehetséges alternatív megközelítések feltárása keresési problémaként azonosítandó, ahol az optimális felé vezető út megtalálása a kitűzött cél.

2.3.6. Az alfejezet összefoglalása

A saját kutatási téma vonatkozásában, a mesterséges intelligenciáról folytatott szakirodalomkutatás az alábbi pontokban járult hozzá célkitűzéseimhez:

- A mesterséges intelligencia fogalomalkotás kereteiben alapvetően az alábbi kritériumok mentén érdemes a modellezési gyakorlatokat elvégezni:
 - vonatkoztasson el az emberi szubjektivitástól (működjön minél inkább kontextus független módon);
 - törekedjen minél komplexebb célfüggvény esetén az optimális megoldás megtalálására a rendelkezésre álló erőforrások függvényében;
 - az emberi üzemszerű észlelőképesség/hermeneutikai meghaladásával, legyen képes a rejtett mintázok felismerésére;
 - maximalizálja a kinyerhető információt az elérhető adatokból;
 - a kinyert információ értékességi aspektusai és hibái hassanak vissza a következő hasonló problémamegoldási folyamatra.
- A modellezést, explicit szabályalkotás nélkül, tehát gépi tanulással érdemes elvégezni a fejezetben ismertetett fejlesztési fázisok követelményeivel összhangban;
- Álláspontom, hogy újszerű innovatív algoritmusok és technikák kutatása és fejlesztése a célra vezető, egy a dolgozatban is ismertetett adatvagyonból történő maximális potenciál kiaknázása érdekében, melyben szerepet játszik a gépi tanuló rendszerek kimeneteinek minősítése azok érték-irány levezetésével.

2.4. Gyanúgenerálás az információbiztonság kutatási területén

2.4.1. A gyanúgenerálás fogalmi kerete

A gyanúgenerálás, mint tudományos szakkifejezés, relatíve ritkán jelenik meg a szakirodalomban. Az elektronikus formában elérhető Magyar Nyelv Értelmező Szótára (Magyar Elektronikus Könyvtár 2016) alapján a gyanú „*általában bizonyos tényekre, körülményekre alapozott sejtlem, felvetés, rendszerint arra vonatkozik, hogy valami rossz, kellemetlen dolog történt vagy fog bekövetkezni*”. A Wikiszótár (2020) alapján a „generál” ige matematikai értelemben is felhasználható, jelentése: „*matematikai eljárással létrehoz vagy terem*”.

A szóösszetétel ezek alapján jelentheti azt a matematikai eljárással levezetett felvetést, amely szerint valamivel (egy objektummal) megvalósult vagy megvalósul egy nem kívánatos tevékenység. Egy adott „*körülmenyre vonatkozó sejtlem, felvetés*” magában hordozza a bizonytalanság fogalmát, azaz a gyanú lehet egy intuíción alapuló ösztön, mely korábbi tapasztalatok mintájára sugallja a gyanú felvetője számára egy negatív esemény bekövetkezését. A gyanúgenerálás, tehát kockázatértékelésként is értelmezhető, ahol a mindenkori cél a gyanús esemény, vagy gyanúmomentum valószínűségét objektíven, a lehető legpontosabban megközelíteni, a becslést ténnyé formálni. Magas szintű párhuzam vonható az anomália észleléssel, mely célja, DUA és DU (2011) értelmezésében, a „*jelentősen eltérő viselkedések detektálása, egy előre meghatározott normális mintától való különbözőség azonosítása*”.

Az anomália, az idézett megfogalmazásból adódóan jelenthet pozitív megkülönböztetést is, mivel a definíció nem tiltja azt, hogy a normális minta alatt elutasító magatartást értsünk, azaz, egy kitűnő diák is lehet anomália a bukottak között. Továbbá, abban az esetben, ha a norma kitűnő diáknak lenni, akkor mindenki más, az előző példa alapján, anomáliának nevezhető. Ugyanerre a gondolatmenetre terel HAN és KAMBER (2011), akik a „*szélsőséges értékek*” fogalma alatt tekintettek olyan elemekre, melyek „*nem felelnek meg az általános viselkedésnek*” és „*durván eltérnek az adathalmaz többi részétől*”, ezért a szélsőséges érték és anomália fogalmát szinonimaként értelmezem.

A magyar nyelvet elemezve a „gyanúsán jó” és „gyanúsán rossz” kijelentéseket érdemes kiértékelni. Ha egy láda almából kiragadva egyet azt mondjuk: ez az alma gyanúsán rossz, akkor feltehetően azt értjük alatta, hogy az alma nem ehető, éretlen, kukacos stb., tehát lemondóan, negatív értelemben nyilatkozunk róla. Amennyiben, a kiválasztott alma esetén azt feleljük: ez az alma gyanúsán jó, akkor kétely merül fel bennünk, pesszimisták vagyunk az alma jóságával kapcsolatban, megkérdőjelezzük annak megfelelőségét. Következésképpen, a gyanús alma jóságát mindkét nézőpontból becsmételjük. A gyanú, ebből kifolyólag, negatív árnyalatot fest, így az anomália észlelése magában foglalja a gyanúgenerálást, azonban gyanúgenerálás alatt, véleményem szerint, a mindenkori rosszat keressük. Azonban az, hogy mi számít rossznak az ideológiánként eltérő lehet. A lopás, vélhetően kivétel nélkül, minden jogi rendszerben elítélendő, mindamellett, Robin Hood a legtöbb ember számára egy pozitív karaktert testesít meg.

PITLIK (2013) kiemeli, hogy „*a gyanúgenerálás lényege, hogy sok dimenzió mentén egy egységes gyanúerőteret legyünk képesek felépíteni, lényegében context free módon*”. Egy osztályban nem kizárólag a jegyek átlaga írja le a diákokat, hanem pl. a szemük színe, magasságuk, sporteredményeik stb., ezért szélsőséges érték nem csak a kitűnő diák lehet, hanem az egyedüli kékszemű, a legmagasabb, vagy a karate aranyérmes. Az, hogy mi számít gyanúsaknak, az csak a döntéshozó szubjektív értékítélete (vö. 2.3.2. alfejezet irány-preferenciák következménye).

PITLIK (2013), továbbá hangsúlyozza, hogy a „*gyanú fogalma az anti-diszkriminációs számításokhoz kell, hogy kötődjön, melyekben a vizsgált jelenségek elsődlegesen semmilyen*

vonatkozása nem kaphat szerepet, ami ezek tartalmát, lényegét, jelentését, egymással való fogalmi kapcsolatát érinti.” Az anti-diszkriminációs matematika filozófiája a „minden objektum másképp egyforma” elv érvényre juttatása (vö. MÉRŐ LÁSZLÓ: Mindenki másképp egyforma című könyve (2007)), azaz csak azt az objektumot tekinthetjük szélsőséges értéknek, amelyik a többdimenziós térben az objektumokat leíró attribútumok közül a legtöbb esetben is különbözik számottevően a többitől.

A fenti példa alapján, anti-diszkriminatív módon, ha létezik kitűnő diák, akinek a szeme a legkékebb, a legmagasabb és a legjobb karate versenyző, ő tekinthető anomáliának (abban az esetben, ha a véges sok leíró ismervv közül már nem tudunk többet felsorakoztatni, ami gyakorlatilag lehetetlen. Azonban, elméletben az anti-diszkriminatív filozófia alapján, véleményem és a kvantumfizika elvárásai szerint, nem is létezhet az univerzumban anomália, mivel minden objektum másképp egyforma), ami megfelel az elméleti fizika azon vélelmének, hogy információ nem veszhet el (vö. az információmegmaradás törvénye fekete lyukak esetén is (HAWKING 2017)). Évéggett, a tanulási halmaz optimalizációjára vonatkozó korábban azonosított problémát anti-diszkriminatív módon pl. hasonlóságelemzéssel érdemes lehet részleteiben kivizsgálni.

A fentiekből következtetve, a dolgozatban a kontrollhiányosságok detektálását gyanúgenerálásnak nevezem, ahol a mindenkorai cél az információbiztonsági kontrollok megfelelőségének prediktálása, a hiányosságok, mint gyanúmomentumok rögzítése. Az automatizált gyanúgenerálás a döntéstámogató rendszer (robot-auditor) alapvető kompetenciája.

2.4.2. Gyanúgenerálás a kibervédelemben

Az információbiztonsági kihívások megoldása gépi tanuló módszerekkel látszólag széles szakirodalmat kínál: a Scopus adatbázisban jegyezett publikációk száma csak az előző 7 évben 5,006 db volt, ahol az “information security” és “machine learning” kulcsszavak konjunktív kapcsolatban kerültek felvitelre a keresőbe (SCOPUS 2021b). Ezt kiegészítve az “audit” szóval, a keresés eredménye kizárólag 32 db, amelyből az összes cikk a sérülékenységelemzés kontextusában használja az audit szót. Ez azt jelenti, hogy feltételezhetően valamennyi publikáció a kiberbiztonságra összpontosít (pl. hálózati forgalom elemző alkalmazások, behatolásérzékelés, stb.), mely az információbiztonság egy kitüntetett szakmai alterülete.

Kevésbé technikaibb szempontból, mely a dolgozatban felhasznált adatvagyonra is jellemző, jelenleg korlátozott számban elérhető, közvetlen információbiztonságra vonatkozó kutatás a Scopus-on kívüli teret vizsgálva, mely azt jelenti, hogy a folyamat-szintű, menedzsment és informatika együttes szabályozási területe (kontrollfolyamatok) jelenleg kevésbé kutatott téma, vagy publikusan nem áll rendelkezésre olyan anyag, mely annak gépi tanulásával történő feldolgozását részletezné.

Kontrollhiányosságokra vonatkozó naplófájlok alapján történő gyanúgenerálással kapcsolatos publikációt nem találtam. Az auditok fogalomkörét vizsgálva, a tudományos társadalom javarészt a gépi tanuló rendszerek auditját kutatja pl. BÜCKER et al. (2020), DAI et al. (2020) és PANIGUTTI et al. (2020), de a kutatási terület felfutása leginkább az utolsó 1 évre jellemző, így a dolgozat a témában történő kutatások egyik úttörőjének is tekinthető.

Érdemes azonban a kibervédelem és gyanúgenerálás közös szakirodalmának és főbb mérföldköveinek az áttekintése, mivel a saját kutatási téma vonatkozásában az hozzáadott értékkel jár a technikai háttér és biztonsági mechanizmusok mögött meghúzódó műszaki tartalom

egyezősége végett (pl. a dolgozatban használt adatvagyon a kriptográfiai folyamatok összefüggéseit is magában hordozza).

Párhuzamot keresve és vonva a kibervédelem területének kutatásaiban, DUA és DU (2011) a gyanúgenerálás fogalmkörébe illeszthető, azzal közeli kapcsolatban álló gépi tanuló megoldásokat két részre tagolja:

1. Az első az „anómália észlelés”, mely elsődleges célja a normálistól való eltérés detektálása, azaz egy új objektum a historikusan begyűjtött adathalmazhoz való összevetése és annak vizsgálata, hogy az mennyire különbözik vagy illeszkedik az eddig definiált mintázatba. A technika, így a jelentősen eltérő viselkedéssel bíró objektumok felderítését állítja középpontba. Erre példa, többek között, a hálózati forgalom monitorozása, ahol a hálózati viselkedés naplófájljait alapul véve a cél a gyanús magatartás detektálása, azonban, az eltérő viselkedés még nem jelent azonnali anomáliát, ezért az anomália észlelés egyik legnagyobb hátránya, hogy az észlelt eltéréseket azonnal ki kell vizsgálni, így ezen rendszerek esetén általában magas a „hamis pozitív” találat. A másik jelen kutatói probléma az anomália észleléssel kapcsolatban, hogy a védekező rendszerek elterjedésével, azok működését a támadók is el kezdték tanulmányozni, így a kiberbűnözők igyekeznek a normális felhasználói magatartást lemásolni, amit a rendszer érvényesnek kategorizálhat, így az anomália tényét nem deklaráló naplóállományok kivizsgálása is indokolt lehet. Továbbá, „hamis negatív” találat alatt, azt a találatot értjük, melyet a rendszer normális tevékenységként könyvelt el, azonban az anomália volt, ezért a „hamis negatív” találatok is magas kockázatokat hordoznak magukban. Összefoglalva, az anomália észlelő/érzékelő rendszerek esetén jelen legnagyobb kihívás a „hamis pozitív” és „hamis negatív” találatok minimalizálása, melyet a dolgozatban ismertetett modellek kiértékelésénél is kiemelten kezeltek.
2. A második csoport a „visszaélés felismerés” vagy „aláírás felismerés”, ahol a cél egy adott új viselkedésről eldönteni, hogy az kártékony-e vagy sem egy meglévő adatbázis alapján, amelyben a gyanús viselkedések definiálva vannak, és az algoritmusok kereséssel győződnek meg arról, hogy az adatbázisban tárolt mintázatba az új jel mennyire illeszkedik. Ezen az elven működik a legtöbb vírusirtó, melyek az ismert vírusok karakterisztikáit letárolják, majd a fájlok szkennelésekor ezt a definíciós adatbázishoz hasonlítják. A legnagyobb hátránya a kialakított módszereknek, hogy bár hatásosnak tűnnek az ismert tulajdonságokkal rendelkező támadásokkal szemben, egy új, eddig nem látott viselkedés esetén nem képesek annak azonnali kiszűrésére.

Kutatói munkám a gyanúgenerálás területével foglalkozik, mely a leírtak alapján az anomália észlelés kategóriájának feleltethető meg, így a kettőt a továbbiakban szinonimaként használom.

Szakirodalmi kutatásom alapján, a gyanúgenerálás kiberbiztonsági alkalmazásai 1987-től kezdtek teret hódítani DENNING (1987) behatolásérzékelés definícióját követően. Mivel a gépi tanulás módszerei akkoriban kevésbé terjedtek el széles körben, többek között köszönhető annak, hogy a gépi tanulásban alkalmazott algoritmusok általánosságban magas hardver performancia igénytel rendelkeznek, a kezdeti kutatások főleg a matematikai statisztika módszereivel igyekeztek a gyanús események felderítésében.

- SMAHA (1988) rendszerhívások (az operációsrendszer és futtatott szoftverek közötti kommunikáció eszköze) naplófájljaira vonatkozó elemzési keretrendszert javasolt, az információbiztonságra irányuló kutatások egyik legelső elméleti úttörőjének tekinthető.
- Hasonlóan SMAHA (1998) kutatásához, sok más kutató is a rendelkezésre álló rendszerhívásokból indult ki, hogy a hálózat, illetve az üzemeltetett szoftverek elleni

támadásokat észlelje. Ide tartozik, többek között, GHOSH et al. (1998), YE et al. (2001) és LIAO és VEMURI (2002) munkái, melyek a matematikai statisztika eszköztárát alkalmazták a gyanúgenerálás módszereként úgyszintén.

- WARRENDER et al. (1999) különböző behatolásérzékelő modelleket épített, szabályalapú algoritmusok és Rejtett Markov Lánckok felhasználásával, mely szakirodalmi kutatásom alapján az első olyan kísérlet volt, ahol a gépi tanulás ideológiájához közelebb álló algoritmusok ideálisabb teljesítményhez vezettek a hagyományos statisztikai eljárásoknál.

A 90-es évek végére a többváltozós statisztikai módszerek közé beékelődtek a gépi tanuló technikák, melyek fokozatosan egyre nagyobb teret nyertek és kezdték leváltani a gyanúgenerálásban a statisztikai modellezés eszközeivel ellátott módszertanokat, így racionális a saját kutatás aspektusából az a döntés, hogy a kontrollhiányosságok detektálására mellőzöm a tisztán statisztikai módszereket a megoldásokra irányuló evolúciós fejlődés eredménytermékeként.

- Első kiemelt szereplője ezen paradigma-váltásnak LEE és STOLFO (2000), historikusan gyűjtött hálózati forgalmat tanulmányoztak és egy adatbányászati keretrendszert javasoltak az anomália detektálásával foglalkozó kutatók részére.
- PORTNOY et al. (2001) klaszterező eljárással kísérelték megállapítani audit fájlokból a betörés tényét és sajátosságát hálózati forgalmat vizsgálva, viszonylag alacsony 50%-os pontossági mutatóval, mely a modell teljesítőképességét tesztadaton mért igazolt találatok számával határozták meg. A kutatás rávilágított, hogy a vizsgált területen, egyrészt, önmagukban a felügyelet nélküli tanulási módszerek nem tűntek elégségesnek, továbbá, az egyszerű tradicionális modell önmagában nem volt elegendő a probléma megoldására.
- YAMANISHI és TAKEUCHI (2001) hibrid módszert alkalmaztak felügyelt és felügyelet nélküli tanulási módszerek közös alkalmazásával, mely publikációjukban az anomália, mint klasszifikációs probléma, pontozásos rendszerrel került kiértékelésre. A klasszifikációt, az adathalmazban meghúzódó mintázat alapján automatikusan feltárt logikai szabályok mentén értelmezték, azaz, logikai következtetések révén vezették le az anomália meglétét a rendelkezésre álló rendszernaplók attribútumai alapján. Átlagosan 71%-ban volt képes a modell az adathalmaz osztályainak pontos meghatározására. A felügyelet nélküli gépi tanuló eljárások további kutatásokban is előfordulnak.
- Például ESKIN et al. (2002) kifejezetten a felügyelt és felügyelet nélküli módszerek hibridizálását ajánlja. Álláspontja szerint a felhasználói szokások gyakran változnak, mely újabb mintákat generál a gyarapodó naplóállományba, így a tanuló eljárások karakterisztikáit és összetételét is dinamikusan változtatni kell, különben magas hamis pozitív találati aránnyal szükséges számolni. ESKIN et al. (2002) szerint felügyelet nélküli módszerekkel lehetséges a felhasználói viselkedések csoportosítása idősorosan is a változások összhangjában.
- YE et al. (2002) behatolásérzékelő modellt terveztek, mely 4 napig folyamatosan gyűjtött naplófájlokat dolgozott fel és csupán 0.42% volt az anomália szerepe. Ez magas teljesítményhez, de magas hamis pozitív találatokhoz vezetett. A publikáció általam leszűrt konklúziója, hogy a magas találati arány még nem jelenti a rendszer jóságát, ahhoz szofisztikáltabb metrikákat kell kialakítani fókuszálva a kiegyensúlyozatlansági osztályeloszlás jellegzetességeire.
- MAHONEY és CHAN (2003) hálózati forgalmat elemeztek idősoros adatokon LERAD szabály alapú tanuló algoritmussal, s alacsony 50%-os teljesítményt értek el, mely az elkészített modell klasszifikációs erejét mérte (A LERAD célja, hogy olyan feltételes szabályokat találjon, amely képes váratlan eseményeket idősorosan azonosítani). Véleményem szerint, a kutatásból levont eredményeket mérlegelve, a szabály alapú algoritmus nem

bizonyult hatékonynak, így annak más módszerekkel való hibridizálása magasabb teljesítményhez vezethet.

- FEINSTEIN et al. (2003) statisztikai módszereket alkalmaztak hálózati forgalom elemzésére, mely célja a túlterheléses támadások kiszűrése. Az erre épített modell magas megbízhatósággal működött és képes volt a gyanú tényét azonosítani, ez azonban annak is köszönhető, hogy túlterheléses támadások esetén a cél az adott hálózat támadása, így folyamatos terhelések miatt a forrás IP címe ismertté válik, mely gyakoribb kéréseket intéz a céleszköz felé, így ez gyorsan észrevehető és blokkolható.
- LEUNG és LECKIE (2005) klaszterező eljárással elemezték hálózati forgalmat, ami gyengébbnek bizonyult, mint a felügyelt tanulási eszközök, SVM-mel 94.9%-os, KNN algoritmussal 89.5%, azonban egy módosított (hibridizált) klaszterező eljárással 97.3%-os teljesítményt értek el az anomália előrejelzésének pontosságára. Álláspontom szerint a kutatási eredmény létjogosultságot teremt a hibrid rendszerek fejlesztésének részletesebb kutatását illetően.
- ZHANG és ZULKERNINE (2006a) a gyanúgenerálás és aláírás felismerés komponenseit hibridizálták, azaz az első olyan példa, mely nem algoritmusokat, hanem magát a modell-logikákat ötvözte. A modell Véletlen Erdő alapú gépi tanuló eljárással kísérte meg az anomália észlelését, és magasabb performanciát voltak képesek a modellel elérni, amit elsősorban az aláírás felismerés kiszűr, majd az anomália észlelő komponens magasabb pontossággal elemmez.
- ZHANG és ZULKERNINE (2006b) egy másik kísérletükben különböző algoritmusokat alkalmaztak behatolás érzékelésre, melyek kifejezetten alacsony teljesítményt produkáltak. SVM-mel 67%-os, míg KNN-nel csupán 11%-os eredményt értek el, melyből triviálisan következik, hogy az algoritmusok önmagukban nem voltak képesek a feladatok megoldására.
- BHUYAN et al. (2011) összegyűjtötték koruk behatolásérzékelő rendszereinek megoldásait, és arra jutottak, hogy az alkalmazni kívánt modellek nem hatásosak teljes mértékben az anomáliák felderítésére, továbbá, véleményük szerint, azon inkrementális tanuló megközelítések tűnnek ígéretesnek, melyek az adatbányászat, neurális hálók és küszöbérték alapú elemzés kombinált összetevőivel rendelkeznek.

Az előzőleg felsorolt kutatási eredmények javarészt hálózati forgalom alapon és rendszerhívások által kísérelték meg az anomália észlelését az információrendszerekben, azonban megközelítőleg 2011-2012-től kezdődően a gyanúgenerálás kutatási területei terjeszkedni látszódnak az információbiztonság területén. A közösségi hálók, a mobil eszközök, az apró internetre csatlakoztatható szerkezetek (pl. IoT) elterjedésével megnőtt az igény, és ezzel együtt a tudományos társadalom kutatási kedve, hogy egyéb területeket is vizsgáljon, mely az anomália felderítésével és elemzésével foglalkozik (pl. HORVÁTH et al. 2016).

- RANJAN és SAHOO (2014) módosított K-közép klaszterező eljárást ajánl különböző távoli külső támadások detektálása érdekében, melyben egyes támadás típusokat (szolgáltatásmegtagadás, távoli rendszertámadás, sérülékenységi kihasználás és portszkennelés) különböző modellekkel vizsgál. A javasolt módosított klaszterező algoritmus teljesítménye rendre 96%, 90%, 71% és 70% volt, mely a hagyományos klaszterező eljárásokon felül teljesített.
- YU és PAREKH (2016) kiemelik a felügyelet nélküli gépi tanuló módszerek egyedüli alkalmazásának hiányosságait. Mivel a felügyelet nélküli algoritmusok címkézetlen adatokkal operálnak, ezért azok teljesítményének mérése megkérdőjelezhető, mert nincs előzetes felülvizsgált adat arra vonatkozóan, hogy egy kiugró érték valóban anomália vagy sem.

Kutatásukban Bayesian hibrid modellt alkalmaznak, és arra a következtetésre jutnak, hogy a hibrid modell teljesít a legjobban hálózati adatokban kutató anomáliák észlelése tekintetében.

- FANAEE et al. (2014) IP/TV hálózati anomália felderítésére többdimenziós megoldást javasol, mivel a hagyományos klaszterező és regressziós eljárások, továbbá a faktor-analízis kizárólag kétdimenziós modellek esetén működtethetők. Kiemeli, hogy a hálózati anomáliákat a felhasználó, attribútum, idő mentén érdemes megvizsgálni, mely 3 dimenziós térbe helyezi a detektálás szükségességét.
- ALKASASSBEH (2018) hibridizált neurális háló alapú eljárással vizsgált túlterheléses támadás áldozatául esett hálózatokat, és arra az eredményre jutott, hogy a hibrid modell képes volt 98.4% teljesítést produkálni, úgy, hogy kizárólag „hamis negatív” hibákat vétett a rendszer.
- SHEKHAR és AKOGLU (2019) együttes módszerekkel gyanút generáltak (kiugró értéként kezelve) bizalmas információ védelmére, mint pl. twitter bot detekció, e-mail spam felismerés, közel 90%-os pontosságot elérve az osztályozás előrejelzésére.

Összegezve a kibervédelem területén relevánsnak tűnő kutatási eredményeket, az a konklúzió szűrhető le, hogy a téma műszaki sajátosságai végett a hibrid megközelítések tűntek célravezetőnek, valamint megerősítést nyert az a tény is, hogy a gépi tanuló modellezés alkalmasabb a tradicionális statisztikai megközelítéseknél. Így a dolgozat a hibridizálást mindenképpen kell, hogy vizsgálja és ennek adaptációját, újszerű megoldásainak feltárását fel kell, hogy vállalja.

2.4.3. Az alfejezet összefoglalása

Az információbiztonsági kontrollhiányosságok detektálására irányuló gyanúgenerálás szakirodalmi áttekintése alapján az alábbi megállapítások és vállalások tehetők:

- A gépi tanuláshoz létjogosultsága van a robot-auditor-probléma megoldásában is;
- A jóságmetrikák ideálisabb értéke érhető el gépi tanuláshoz támaszkodva, szemben a statisztikai eljárások alkalmazásával;
- A hibrid megközelítés alkalmasabbnak bizonyult a definiált problémák megoldására, mely hibridizálás kivitelezhető különböző technikák (pl. felügyelt és felügyelet nélküli módszerek) közös felhasználásával, így a hibridizáció hatásait a robot-auditor kapcsán is fel kell tárni.

2.5. Hipotézisek felállítása a szakirodalmi áttekintés alapján

A 2. fejezetben ismertetett, az információbiztonsági auditokra, mesterséges intelligencia fogalomkörre és gyanúgenerálásra vonatkozó szakirodalom mélyreható tanulmányozása alapján a célkitűzésekkel összhangban, az alábbi hipotéziseket állítom fel:

- **H1:** Az információbiztonsági auditjelentések szöveges eredményeiből strukturált adatbázist alkotva és bemenetként a mesterséges intelligencia fogalomkörébe illeszthető eszközökkel azt feldolgozva, az auditok során feltárni kívánt kontrollhiányosságok megléte a véletlen találgatásnál nagyobb valószínűséggel kimutathatók, azaz a kontrollhiányosságok konstellációi matematikailag értelmezhető összefüggéseket hordoznak magukban.
 - **H1.1:** A gyanúgenerálás, mint megoldandó üzletileg értelmezett probléma sajátosságait értékelve, a kontrollhiányosságok detektálása megoldható felügyelt és felügyelet nélküli gépi tanuló eljárásokkal is.
 - **H1.2:** A gyanúgenerálás teljesítménye fokozható hibrid megközelítésben, azaz a felügyelt és nem felügyelt módszerek együttes felhasználásának a kutatásban alkalmazott releváns performancia metrikái ideálisabb értékeket mutatnak, mint önálló alkalmazásban
 - **H1.3:** A hibrid modell többlet-információs értéket teremtve képes az egyszerű modellek általánosító képességén javítani.
- **H2:** A döntéstámogató rendszer genetikai potenciálja letapogatható hasonlóságelemzéssel ellátott kereső eljárással a tanításra alkalmazott adathalmaz irányított feldolgozásán keresztül, úgy, hogy a genetikai potenciálhoz vezető kereső eljárás a genetikai algoritmusok esetén alkalmazott véletlen mutáció és a populáció egyedeinek keresztezése nélkül is képes ideálisabb eredményt szolgáltatni.
- **H3:** A mesterséges intelligenciával ellátott döntéstámogató rendszerek teljesítményalapon a gépi tanuló alkalmazások klasszikus tesztelési eljárásai nélkül is rangsorolhatók, a predikciók, mint generált gyanúforrások leíró tulajdonságainak érték-irány levezetésével és az ezen adatokat feldolgozó matematikai apparátussal, mely automatizáltan képes a preferált modellek objektív meghatározására.

3. ANYAG ÉS MÓDSZERTAN

3.1. Adatgyűjtés

A hipotézisek deklarálását követően azok alátámasztását/elvetését igazoló kísérletek elvégzéséhez az adatbázisok megtervezése és feltöltése volt a feladat, ahol kritikus elemként szolgált olyan adatbázisok reális, megbízható és objektív összeállítása, melyek már képesek a modellezési célkitűzések megvalósítására, azaz tiszta képet adnak arról, hogy melyek azok az információbiztonsági kontrollterületek, melyek a legnagyobb kockázatot jelentik a szervezetek számára a biztonságos üzemeltetés szempontjából. Mivel az információbiztonsági kitétségek, a kontrollok hiánya, és azok nem megfelelő üzemeltetése üzleti titok és a legtöbb szervezet számára sérülékenységet jelent azok nyilvánosságra hozatala, ezért a szervezetek a legtöbb esetben nem hajlandók azok megfelelőségéről nyilatkozni, vagy eltakarva a teljes igazságot, képesek a hiányosságok ellenkezőjét is jelenteni akár szervezeten belül is félve a fegyelmi eljárások következményeitől. Ezen gyakorlati úton is tapasztalt érvek alapján indokoltnak tűnt az információbiztonsági kontrollok megfelelő működéséről tisztább képet kapni aggregált szinten független ellenőrző szervek által, akik a szervezetek vezetésének és munkavállalóinak befolyásától mentesen végeznek információbiztonsági vizsgálatokat és jelentik a szervezetek számára a hiányosságokat és megállapításokat egy riport, auditjelentés formájában. Mindazonáltal, mivel humánerőforrás által gyűjtött adatokról van szó, nem pedig teljesen objektív fizikai mérésekről, így továbbra is fennáll az adatgyűjtésből eredő pontosság és teljesség (adatintegritás) kockázata, mely alapvetően a társadalomtudományi jelenségek vizsgálatánál külső adottság.

A kutatáshoz felhasznált adatok forrása egy saját (megtervezett és feltöltött) adatbázis, melyben tartalmazott valós adatanyagot helyszíni vizsgálat (terepmunka) alatt gyűjtöttem két könyvvizsgáló, auditor és tanácsadó szervezettől – ezek hozzájárulása mellett, melyek mindegyike 500 főnél több munkavállalót foglalkoztat.

Az adatbázis a lefolytatott auditok megállapításainak számát tartalmazza témakörönként megbontásban az ISO/IEC 27001:2013 szabvány A melléklet információbiztonsági kontrolljai alapján (összesen 114 + 1 addicionális attribútum), továbbá, magában foglalja az audit típusát, az auditált szervezet iparági besorolását, és az audit hatókörében lévő informatikai rendszerek számát. Az ISO szabványcsalád világszerte elismert, az ISO/IEC 27001:2013 magas színvonalon és könnyen értelmezhető struktúrában fogalmazza meg az információbiztonsági kontrollkövetelményeket, ezért ésszerű volt az adatbázis attribútumainak létrehozása a szabvány által definiált formában.

Az auditor szervezetek összesen 127 valós auditjelentést bocsátottak rendelkezésemre szöveges, anonimizált formában 2016-2020 között lezárult auditokról. Az auditjelentések 127 különböző egymástól független auditot reprezentálnak. A legrövidebb auditjelentés 5 oldalas, míg a leghosszabb 106 oldalas volt. Az adatok gyűjtésére 2020. április 1. és 2020. június 30. között került sor és összesen 3,699 oldalnyi auditjelentés került feldolgozásra (átlagosan kerekítve 29 oldal auditjelentésenként).

Szükséges megjegyezni, hogy a disszertáció készítése során tudatosan kerültek kizárásra kérdőívre és/vagy interjúra alapozó adatgyűjtések, mivel egyrészt, a kontrollhiányosságok meglétét egy vezető nem szándékozik külső féllel megosztani, másrészt, ebben a témakörben különösen csak a függetlenül gyűjtött adat az, ami reprodukálható, objektíven ellenőrizhető módon írja le a valóságot.

Még a fenti adatbázis kapcsán is fennáll annak esélye, hogy adott auditjelentés feltárni vélt hiányosságai csak az eljáró személyek számára tűntek hiányosságnak, míg más auditorok ezeket

esetleg nem vették volna figyelembe. Ez az objektivitási korlátozás a maximum, amit egy disszertáció kapcsán akceptálni lehet valós probléma valós adatvagyonára estén, ennél több adatminőségi kockázat már eleve önkényessé tenné a hipotézisek értékelését.

A terepmunka során összesen 5 különböző audit típust lehetett elkülöníteni:

- **Könyvvizsgálathoz kapcsolódó informatikai vizsgálat:** A könyvvizsgálat elsődleges célja, hogy a könyvvizsgáló megbizonyosodjon arról, hogy a könyvvizsgált fél a vagyoni helyzetéről megbízható pénzügyi adatokat publikált, azaz a pénzügyi beszámoló teljessége és pontossága biztosított. Ezt a könyvvizsgáló különböző pénzügyi és számviteli tesztek, könyvelési folyamatok és kapcsolódó bizonylatok és számlák felülvizsgálata révén ellenőrzi. A könyvvizsgálat szerves részét képezi az információbiztonsági audit, azaz a könyvvizsgálók informatikai szakértők bevonásával ésszerű bizonyosságot szándékoznak szerezni, hogy a könyvvizsgált fél által a pénzügyi beszámolóra hatással lévő informatikai rendszerek üzemeltetése kontrollált körülmények között zajlik. Mivel papír-alapú könyvelési gyakorlatok csak elvétve fordulnak elő, ezért az informatikai rendszereket is tüzetesen át kell vizsgálni, hogy a könyvvizsgáló meggyőződjön arról, hogy a tárolt, továbbított, feldolgozott pénzügyi adatokhoz nem történt illetéktelen hozzáférés, adatmódosítás, vagy egy esetleges rendszerhiba nem okozta azok integritásának és mindenkor rendelkezésre állásának sérülését.
- **Jogszabályi megfeleléségi vizsgálat:** A jogszabályi auditok célja egy adott rendelet, jogszabály, direktíva vagy jogi ajánlás által támasztott követelményrendszernek való megfelelés vizsgálata. A terepmunkán az alábbi jogszabályok voltak az auditok hatókörében a teljesség igénye nélkül:
 - 19/2017. (VII. 19.) MNB rendelet a pénzmosás és a terrorizmus finanszírozása megelőzéséről és megakadályozásáról szóló törvény végrehajtásának az MNB által felügyelt szolgáltatókra vonatkozó, valamint az Európai Unió és az ENSZ Biztonsági Tanácsa által elrendelt pénzügyi és vagyoni korlátozó intézkedések végrehajtásáról szóló törvény szerinti szűrőrendszer kidolgozásának és működtetése minimumkövetelményeinek részletes szabályairól.
 - 42/2015. (III. 12.) Korm. rendelet a pénzügyi intézmények, a biztosítók és a viszontbiztosítók, továbbá a befektetési vállalkozások és az árutőzsdei szolgáltatók informatikai rendszerének védelméről.
 - 45/2018. (XII. 17.) MNB rendelet a pénzmosás és a terrorizmus finanszírozása megelőzéséről és megakadályozásáról szóló törvény végrehajtásának az MNB által felügyelt szolgáltatókra vonatkozó, valamint az Európai Unió és az ENSZ Biztonsági Tanácsa által elrendelt pénzügyi és vagyoni korlátozó intézkedések végrehajtásáról szóló törvény szerinti szűrőrendszer kidolgozásának és működtetése minimumkövetelményeinek részletes szabályairól.
 - A Magyar Nemzeti Bank 2/2017. (I.12.) számú ajánlása a közösségi és publikus felhőszolgáltatások igénybevételéről.
 - A Magyar Nemzeti Bank 4/2019. (IV.1.) számú ajánlása a közösségi és publikus felhőszolgáltatások igénybevételéről.
 - AZ EURÓPAI PARLAMENT ÉS A TANÁCS (EU) 2016/679 RENDELETE (2016. április 27.) a természetes személyeknek a személyes adatok kezelése tekintetében történő védelméről és az ilyen adatok szabad áramlásáról, valamint a 95/46/EK rendelet hatályon kívül helyezéséről (Általános Adatvédelmi Rendelet - GDPR).
 - Vezetői körlevél az elektronikus úton megkötött írásbeli szerződésekről, megtett írásbeli jognyilatkozatokról.

- **Szolgáltató Szervezetek Tanúsítása:** Harmadik felek bizonyosságot nyújtó auditja, ahol a harmadik felek egy informatikai szolgáltatást nyújtanak ügyfelek számára (pl. informatikai üzemeltetés, fejlesztés, adatfeldolgozás, stb.) és az ügyfelek igénylik a szolgáltató szervezet auditálását annak meggyőződésére, hogy annak információbiztonsági kontrollkörnyezete a szerződésekben rögzített módokon történik. Az auditot SOC1 és SOC2/3 vizsgálatoknak rövidítik (Service Organization Control), ahol a SOC1/2/3 az audit típusát jelenti, és az ISAE 3402 és ISAE 3000 audit standardok alapján történik a vizsgálat (BARTA 2020).
- **ISO/IEC 27001:2013 vizsgálat/réselemzés:** Az ISO/IEC 27001:2013 szabvány követelményei szerinti réselemzés, valamint a tanúsításra felkészítő auditok.
- **Egyéb vizsgálatok:** Azon auditok, melyeket nem lehet besorolni az előző kategóriákba. Általában olyan vizsgálatok tartoznak ide, melyek egyedi megbízásokon alapulnak, ahol a vezetőség kérése egy adott rendszer, folyamat, leányvállalat auditálása volt a belső szabályozói környezet vagy iparági jógyakorlatoknak való megfelelés ellenőrzésére, testreszabott hatókörrel.

Az adatbázis nyers attribútumai az ISO/IEC 27001:2013 A melléklete alapján kerültek meghatározásra, melynek kontrollterületenkénti összefoglalását a 2.2.5. alfejezetben ismertetett 2. táblázat, valamint azok részletesebb kifejtését a 2. számú melléklet tartalmazza.

A saját adatbázis 1 további kontrollkövetelménnyel kiegészítésre került, melyet nem tartalmaz a szabvány, azonban számos esetben a vizsgált auditok hatókörét képezte: IT kockázatkezelésre vonatkozó kontrollok, azaz annak tervezése, elkészítése, teljessége, pontossága és felülvizsgálata, illetve kapcsolódó folyamatok összessége, melyet a dolgozatban az A19-es azonosítóval jelölök.

Az adatbázis következetes felépítéséhez az alábbi feltételeket és kritériumokat rendeltem a terepmunkán megismert auditjelentések alapján:

- Az informatikai rendszerek számánál kizárólag az auditok hatókörében szereplő alkalmazások száma került feltüntetésre. Egy alkalmazáshoz, azonban, tartozhatott több, egyéb informatikai elem is mint pl. az alkalmazást kiszolgáló adatbázis-kezelő rendszer, vagy szerveroldali operációsrendszer. Amennyiben vonatkozott megállapítás egyéb infrastruktúra elemre is az alkalmazással összefüggésben, a megállapítás az adatbázisba bejegyzésre került.
- Több esetben, főleg a könyvvizsgálatokhoz kapcsolódó informatikai vizsgálatok esetén, az auditok kitértek több szintre is, azaz tesztelve lettek a hatókörben lévő kontrollok tervezet, implementáció és működési hatékonyság szintjén is, melyek egymásra épülnek. Ha a tesztelt szintek valamelyikén volt megállapítás, az bejegyzésre került.
- Egy audit-megállapítás alapvetően egy bejegyzést kapott az adatbázisban a megfelelő kontroll attribútumában. Amennyiben 1 audit megállapítás két kontrollterületet is érintett, mind a két területhez történt bejegyzés. Például, nem megfelelő konfigurációs beállítások és azok naplózásának a hiánya, bár 1 megállapításban szerepelt, két kontrollterületet foglal magába.
- Kizárólag az információbiztonsági kontrollok kerültek feljegyzésre. Amennyiben az auditjelentés kitért üzleti és jogi folyamatok tesztelésére is (pl. GDPR megfelelési vizsgálat), azok nem kerültek az adatbázisba bejegyzésre.

- Bizonyos jelentésekben az auditált félnek lehetősége volt vezetői válasz biztosítására, és olyan eset is volt, melyben az auditált fél nem értett egyet a megállapítással. Ezekben az esetekben a megállapítás bejegyzésre került függetlenül az auditált véleményétől.
- Egyes auditjelentések egy egész cégcsoportra vonatkoztak, ahol pl. egy központi integrált rendszer volt az audit hatóköre. Ez kizárólag egy rekordot jelentett az adatbázisban, nem került bejegyzésre minden egyes tagvállalat a redundáns adatbevitelt elkerülve.

3.2. Alkalmazott algoritmusok és statisztikai eljárások

Az alábbi alfejezet szolgáltatja a hipotézisek bizonyításához felhasznált modellezési gyakorlatokat, melyek bemenetét az előző alfejezetben tárgyalt adatvagyon, valamint kapcsolódó számítási eredmények képezték.

A dolgozatban ismertetett gépi tanuló algoritmusokról és statisztikai eljárásokról részletesebben az alábbi, 4. táblázatban közölt források mélyebb tanulmányozásra adnak lehetőséget a teljesség igénye nélkül.

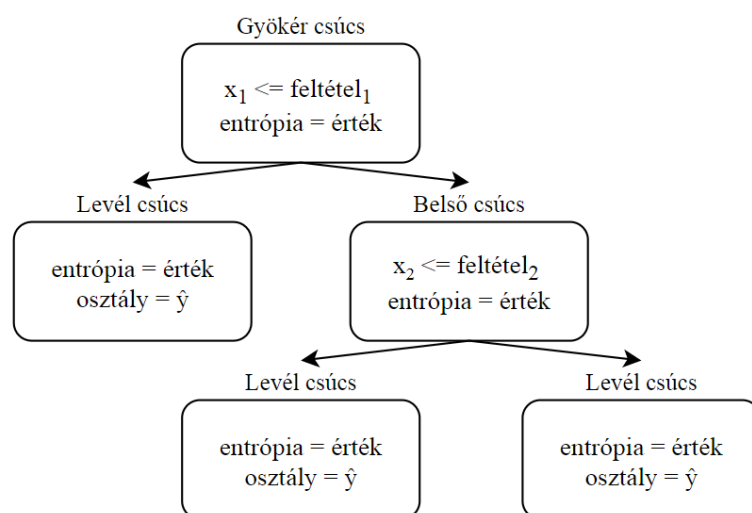
4. táblázat: A dolgozatban ismertetett modellezési gyakorlatok szelektált forráshivatkozásai

Modellezési gyakorlatok	Forráshivatkozások időrendben
<i>Gépi tanulásról általánosságban</i>	HACKELING 2014 RASCHKA 2015 HEARTY 2016 RASCHKA – MIRJALILI 2017 GÉRON 2017 MCCLURE 2017 BURKOV 2019
<i>Döntési fa</i>	IZZA et al. 2020 GELBOWITZ 2021
<i>Neurális háló</i>	BORGULYA 1998 HORVÁTH 2006 GOODFELLOW et al. 2016 CHOLLET 2018
<i>Adaptív Boosting</i>	FREUND – SCHAPIRE 1996 DARÓCZY et al. 2021
<i>Gradiens Boosting</i>	MASON et al. 1999 FRIEDMAN 2001 WADE 2020
<i>Hasonlóságelemzés</i>	DOBÓ 1992 BÁNKUTI 2010 PETŐ 2013 PITLIK 2014

Forrás: Saját szerkesztés

3.2.1. Döntési fa

A döntési fa alapú eljárások az adathalmaz egy kitüntetett attribútumából (gyöker csúcs) kiindulva kérdések sorozatát fogalmazza meg, ahol mindegyik soron követő kérdés (belső csúcs) egy adott attribútum küszöbértékére (feltétel) vonatkozik. Az osztályozandó rekord értékei alapján a döntési fa kijelöl egy útvonalat a rekord számára, mely végül a döntési levél csúcsa által leírt osztályba fogja azt sorolni. Az útvonalat meghatározó feltételek mentén, a döntési fa a bejárható útvonalat, a legtöbb alkalmazott esetben ketté (a kombinatorikai tér csökkentése végett, tehát erőforrás-felhasználási megfontolásokból), de akár több részre is vághatja (csomópont). Az attribútumok és feltételek kiválasztása a döntési fa által nem önkényes módon történik, hanem egy előre definiált kritérium szerint, így az algoritmus automatizált szabályalkotásra képes. A modellezésekhez használt kritérium a dolgozatban az entrópia, mely alkalmas az információ-nyereség kifejezésére (PROVOST – FAWCETT 2013). A döntési fa működési logikáját a következő ábra szemlélteti (11. ábra).



11. ábra: A döntési fa működési logikája

Forrás: Saját szerkesztés

Az információ-nyereség (IG) optimalizálására definiáljuk az alábbi célfüggvényt, ahol e jelölje az entrópiát, N az X mátrix rekordjainak számát, b a baloldali, j a jobboldali belső-/levélcsúcsot:

$$IG(X, f) = e(X) - \frac{N_b}{N} e(X_b) - \frac{N_j}{N} e(X_j)$$

Az entrópia definíciója, ahol $p(i|t)$ a minták azon aránya, amely egy adott t csomópont esetében az i osztályba sorolandó:

$$e(t) = - \sum_{i=1}^n p(i|t) \log_2 p(i|t)$$

A döntési fák alapvető előnyei a dolgozat fókuszában álló modellezés aspektusából:

- a rendelkezésre álló adathalmazt kevésbé szükséges előkészíteni, azaz transzformálni, mivel belső döntési mechanizmusai révén érzéketlen a felhasznált adatvagyon attribútumai között fennálló skála-eltérésekre;
- az algoritmus fehér doboz jellegű, amely azonban nem kerül kihasználásra a kutatásban, mivel együttes módszerek alaposztályozóként történik a módszer használatba vétele. Ennek magyarázata, hogy az együttes módszerek a szakirodalmi áttekintőben tárgyaltakkal (2.3.3. alfejezet) összhangban, képesek az alaptanulók jóságát szignifikánsan javítani (ahogy például az DELGADO et al. (2014), BARTA (2018c) és BARTA – PITLIK (2020) kísérletében is tetten érhető);
- a döntési fákat relatív gyors tanítani;
- képesek a nem-lineáris leképezésre;
- nincs szükség a paraméterek mélyreható finomhangolására (összevetve pl. az SVM-mel);
- így megfelelő alaposztályozót jelent együttes módszerekben (GUIDOTTI et al. 2018).

A döntési fák legfőbb hátránya:

- az egyszerű, azaz nem együttes megközelítésben alkalmazott döntési fák egyik megkerülhetetlen hátránya azok instabilitása, melyet az adathalmaz akár minimális változtatása is indukálhat, amit az együttes felhasználás képes mérsékelni (BARTA – PITLIK 2018a).

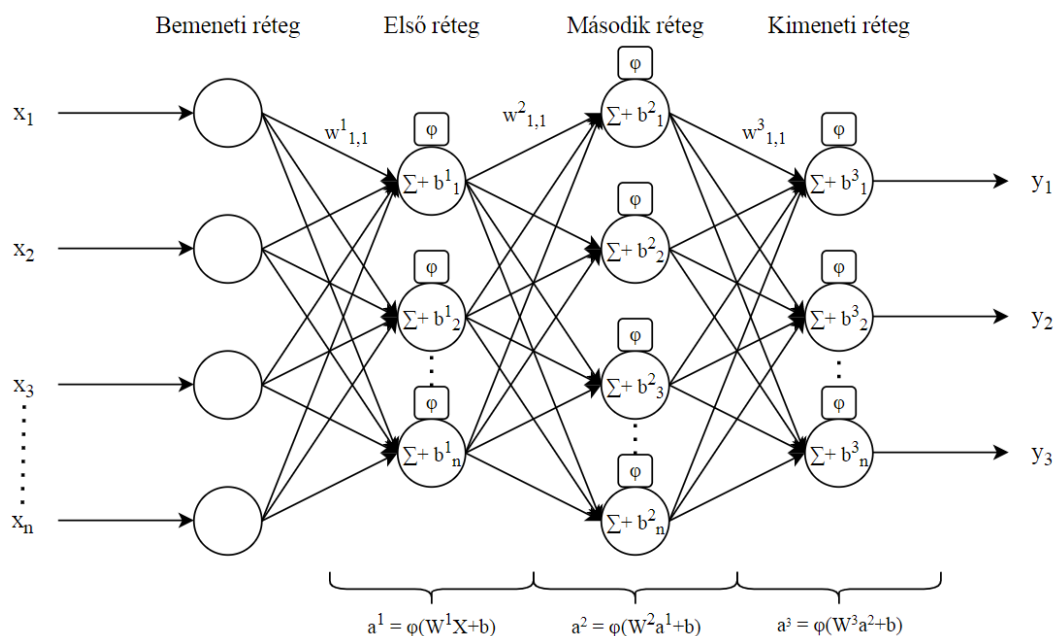
A döntési fa együttes módszerek alaposztályozóként került felhasználásra a döntéstámogató rendszer (robot-auditor) fejlesztésében.

3.2.2. Neurális háló

A neurális háló koncepciója és a mögötte álló problémamegoldó képesség gondolata, több mint 70 éve foglalkoztatja a számításelmélettel foglalkozó szakembereket és kutatókat, kezdve WARREN MCCULLOCH és WALTER PITTS (1943) által megalkotott MCP (McCulloch-Pitts) neuronnal, mely a legelső olyan matematikai modell volt, mely az emberi agy működését szándékozott lemásolni. Ezáltal, kezdetben a klasszifikációs problémákra egy olyan új megközelítést formáltak, mely a korábban látott minták alapján képes volt helyesen csoportokba rendezni az új elemeket. A neurális hálók kutatása, azonban hullámszó volt az elmúlt 50 évben, mivel egyre bonyolultabb algoritmusok láttak napvilágot, és a számítógépek akkori kapacitása nem volt elegendő a nagy mennyiségű számítások elvégzésére, amit a neurális háló modellek produkáltak. 2012-től kezdődően a neurális hálók aranykorukat élik, mely dátumot egyes kutatók a „mesterséges intelligencia forradalmának” neveznek pl. JEFFREY DEAN (2020) a Google kutatója, vagy ANDREW NG (2018) a Coursera alapítója és a „Google Brain” korábbi elnöke. A neurális hálók „mélyített” változatát a szakirodalom „Deep Learning”-nek nevezi, mely a neurális hálók rejtett rétegeinek megnövelt számára utal (GOODFELLOW et al. 2016).

A neurális hálók robbanásszerű népszerűségnek örvendenek, felhasználási területük kiterjed pl. a gépi látás, hang- és beszéd felismerés, irányítás-technológia, stb. kutatási és alkalmazási területeire, a Scopus (2021c) adatbázisa alapján 2020-ban 50916 db, 2019-ben 40866 db, míg 2018-ban 21900 db publikáció volt elérhető, mely tartalmazta a „Deep Learning” kifejezést a címben és/vagy absztraktban. A neurális hálók felhasználhatók a gépi tanuló rendszerek valamennyi különálló megközelítésében, így alkalmazhatók felügyelt, felügyelet nélküli és megerősítéses tanulásban is (LAPAN 2018).

A neurális hálók elméleti megközelítése az emberi agy biológiai neuronjainak gépi formába öltött megtestesítése, amiről a megnevezését is örökölte (KÁSA 2018). Hasonlóan a természetes neuronok felépítéséhez, a mesterséges neuronok is egymáshoz kapcsolódnak, ahol egy közvetítő közeg szállítja a feldolgozott információt neuronról-neuronra. A neuronok különböző rétegekben helyezkednek el: a bemeneti réteg fogadja a feldolgozáshoz szükséges adatokat, a kimeneti réteg szolgáltatja a predikciót, a közbenső rétegek, pedig transzformálják és segítségével jut el az információ a bemeneti neurontól a kimeneti neuronig. A kapcsolóelem a neuronok között a „dendrit”, mely létezhet bármely neuron pár között, de a kapcsolat mértéke a legtöbb esetben az, hogy minden neuron kapcsolódik minden soron következő neuronhoz. A transzferált információ az úton az egyik neurontól a másik felé „átalakul” a súlymátrixok és aktivációs függvények közbenjárásával, melynek összegzett és transzformált értéke a szinaptikus súly. A súlymátrixok a neurális háló „tudásanyagát” tartalmazzák, és a neurális háló a megadott tanulási eljárás és az adatokban fellelhető minta szerint módosítja azt. A neurális háló a definiált veszteségfüggvény szerint kalkulálja a tényadat és becslés közötti különbséget, melyet visszafejt (pl. backpropagation) a háló súlymátrixainak módosítására (gradiens eljárás) (BARTA 2018d, BARTA – PITLIK 2018b). A neurális háló grafikusán ábrázolt modelljét az 12. ábra szemlélteti, mely leírja annak kalkulációs logikáját a bemeneti rétegtől a kimeneti réteig.



12. ábra: A neurális háló működési logikája

Forrás: Saját szerkesztés

A neurális háló aktualizálja minden egyes iterációban a súlyvektorok értékeit:

$$\frac{\partial}{\partial w_j} L(w) = \sum_i (y^i - a^i) x_j^i.$$

A neurális háló legfőbb előnyei, melyek a kutatás szempontjából elsőrendűek és az alkalmazását indokolja:

- képes az univerzális approximációra, (bizonyos feltételek teljesülése mellett pl. nem-lineáris aktivációs függvényt szükséges alkalmazni) (PALUZO-HIDALGO et al. 2020);
- lehetővé teszi a nem-lineáris leképezést (az előző pontból következően);
- rugalmas (felhasználható regressziós és osztályozási problémák megoldására – a dolgozatban kizárólag osztályozóként funkcionál);
- nem rendelkezik a statisztikai módszerek önkorlátozó feltételeivel.

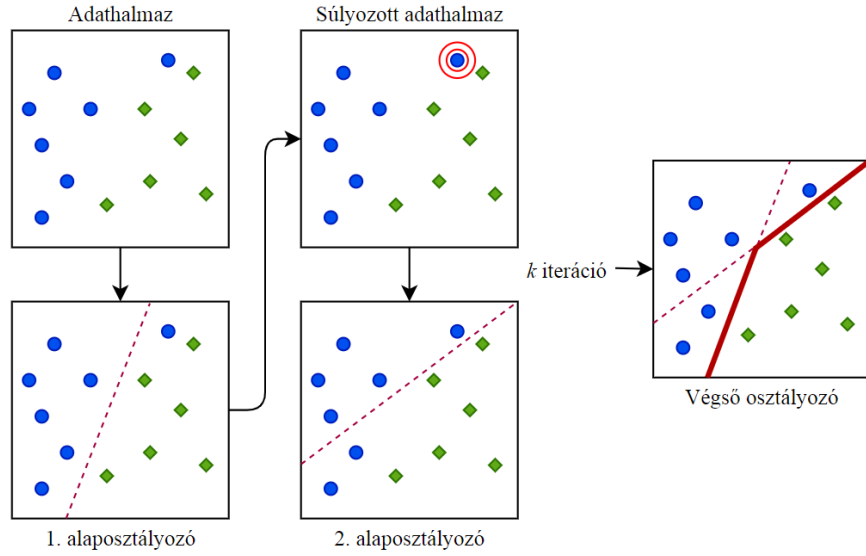
Hátránya, hogy:

- adat- és számításigényes;
- alapvetően fekete doboz modell (bár részlegesen domesztikálhatók (pl. VAUGHAN et al. 2018));
- számos konfigurációs paraméterrel rendelkezik, mely növeli a modellezés folyamatának időszükségletét.

A neurális háló a döntéstámogató rendszer (robot-auditor) fejlesztésében került felhasználásra.

3.2.3. Adaptív Boosting

Az Adaptív Boosting (Adaptív Gyorsítás/Fokozás vagy AdaBoost vagy ABM – Adaptive Boosting Machine) egy együttes módszer, mely a hibásan osztályozott rekordok újrasúlyozásával, több alaposztályozó algoritmus felhasználásával egy erős osztályozót alkot. Az ABM technikát először FREUND és SCHAPIRE (1996) publikálta, mely algoritmusért 2003-ban Gödel díjban részesültek az Európai Elméleti Számítástudományi Egyesület (2003) által. Az ABM egy szekvenciális együttes módszer, tehát az egyes alaptanulók egymásra épülnek (nem függetlenül operálnak), melyben a hangsúly a bemeneti adatokon összpontosul. Minden tanulásra felhasznált adathalmaz meghatározott mértékben az előző alaptanulótól függ, fokozatosan javítva a megelőző osztályozók által elkövetett hibákat. Az ABM az alaptanulók egyes predikcióit kombinálja legvégső osztályozásként, ahol veszi a részeredmények móduszát (többségi szavazás). Alapvetően döntési fáknál alkalmazott módszer, azonban karakterisztikái miatt bármilyen osztályozó szekvenciális összekapcsolására alkalmazható. Az ABM előzőekben leírt tanulási mechanikáját a 13. ábra szemlélteti.



13. ábra: Az ABM működési logikája

Forrás: Saját szerkesztés

Az ABM algoritmus az alábbiakban leírt módon prediktálja a célváltozó értékét:

1. Inicializáljuk a w súlyvektort, ahol:

$$\sum_i w_i = 1.$$

2. Minden c_j alaposztályozóra végezzük el az alábbiakat k iterációban:

- i. Tanítsuk c_j osztályozót a súlyozott adathalmazon:

$$c_j = T(X, y, w).$$

- ii. Prediktáljuk az y célváltozót:

$$\hat{y} = P(c_j, X).$$

- iii. Számítsuk ki a súlyozott hibát:

$$E = w \cdot (\hat{y} \neq y).$$

- iv. Számítsuk ki a koefficiens:

$$\alpha_j = 0.5 \log \frac{1-E}{E}.$$

- v. Számítsuk ki az új súlyvektort:

$$w := w * \exp(\alpha_j * \hat{y} * y).$$

- vi. Normalizáljuk a súlyvektort:

$$w := \frac{w}{\sum_i w_i}.$$

3. Határozzuk meg a végző predikciót:

$$\hat{y} = \left(\sum_{j=1}^k (\alpha * P(c_j, X)) > 0 \right)$$

Az ABM legfőbb előnyei:

- együttes módszer lévén alkalmas a modell pontosságának javítására;
- rugalmasan kombinálható alaposztályozókkal, pl. döntési fákkal;
- alacsony számú konfigurációs paraméter finomhangolása szükséges;
- egyszerűen implementálható.

Az ABM legfőbb hátrányai:

- hajlamos lehet a túlilleszkedésre, mivel a kiugró értékek ismételt súlyozásával kivételkezelést alkot;
- fekete doboz modell.

Az ABM a döntéstámogató rendszer (robot-auditor) fejlesztésében került felhasználásra.

3.2.4. Gradiens Boosting

A Gradiens Boosting (Gradiens Gyorsítás/Fokozás vagy GBM – Gradient Boosting Machine) az ABM algoritmushoz hasonlóan egy szekvenciális együttes módszer, tulajdonképp, az ABM egy finomhangolt variánsa, azonban a differenciálható veszteségfüggvény optimalizálásával (gradiens) törekszik az erős osztályozó algoritmus megalkotására az alábbi módon (MASON et al. 1999, FRIEDMAN 2001):

1. Inicializáljuk a modellt konstansként, ahol:

$$C_0 = \arg \min_c \sum_{i=1}^n L(\hat{y}, y_i).$$

2. k iteráción keresztül végezzük el az alábbiakat:
 - i. Számítsuk ki a veszteségfüggvény negatív gradiensét (pszeudo-reziduális):

$$g_i^k = - \left[\frac{\partial L(\hat{y}, y_i)}{\partial \hat{y}} \right]_{c=C^{k-1}}, i = 1, \dots, n.$$

- ii. Tanítsuk a kiválasztott alaposztályozót a gradiens értékek felhasználásával:

$$c^k := T(X, g_i^k)_{i=1}^n$$

- iii. Aktualizáljuk a modellt a kalkuláltak alapján:

$$C^k = C^{k-1} + \gamma c^k.$$

3. Határozzuk meg a végső predikciót:

$$\hat{y} = C(X)$$

A GBM legfőbb előnyei:

- együttes módszer lévén alkalmas a modell pontosságának javítására;
- rugalmasan kombinálható alaposztályozókkal, pl. döntési fákkal;
- sokféle veszteségfüggvényt képes optimalizálni.

Az GBM legfőbb hátrányai:

- hajlamos lehet a túlilleszkedésre, mely a kiugró értékek által generált hibák minimalizálásából eredhet;
- fekete doboz modell;
- számításigényes.

A GBM a döntéstámogató rendszer (robot-auditor) fejlesztésében került felhasználásra.

3.2.5. Kollaboratív szűrés

A kollaboratív szűrés az ajánlórendszerek esetén egy bevált módszer, mely a hasonlóságot mutató objektumok alapján kísérel meg egy adott objektum számára attribútumot „ajánlani”, mely az ajánlórendszerek gyakorlati példája alapján lehet egy termék két hasonló vásárlási szokással bíró fogyasztó között. A dolgozatban ismertetett kontrollhiányosságra irányuló gyanúgenerálásban az egymással hasonlóságot mutató objektumok az auditjelentések, az ajánlott termékek, pedig a gyanúmomentumok.

A módszer távolság/kapcsolat-metrikákkal képes a hasonló (de itt nem hasonlóságelemzéssel kezelt) auditjelentések beazonosítására, evégett egy speciális klaszterező eljárásról van szó, ahol a leghasonlóbb objektumok egy listával térnek vissza, mely a feltételezett gyanús kontrollokat tartalmazza. A lista a két pl. leghasonlóbbnak ítélt objektum közötti különbségeket adja kimenetként, ahol súlyszámként az objektumok közötti hasonlóságok mértéke felhasználható. A döntéshozó preferenciája, hogy meghatározza az elfogadható hasonlóság küszöbértékét, vagy azon minimum és maximum auditjelentések számát, melyek a legközelebb találhatók a kiválasztott objektumhoz a többdimenziós térben.

A dolgozatban alkalmazott kollaboratív szűrés variáns legfőbb előnyei:

- egyszerűen implementálható;
- más aspektusból közelíti a gyanúgenerálást, így egy konzisztencia-rétegnek tekinthető, mely beépíthető az osztályozó eljárásokba és elősegíti a hibrid modellezést;

A kollaboratív szűrés legfőbb hátrányai:

- A dolgozat szempontjából kiemelendő hátrány, hogy a leghasonlóbb objektumok küszöbértéke önkényes;
- Nehézkes teljesítményének visszamérése.

Kollaboratív szűrés a döntéstámogató rendszer (robot-auditor) fejlesztésében került felhasználásra.

3.2.6. Hasonlóságelemzés

A hasonlóságelemzés egyik eljárása a „mindenki másképp egyforma” elv matematikai megtestesítője, melyben a modellezés célja a legideálisabb objektum megtalálása, oly módon,

hogy definiálunk egy fiktív célváltozót, mely értéke konstans, az objektumokat leíró attribútumok, pedig a megadott irány-preferenciák alapján rangsorolva vannak, azaz idealitás alapján attribútumonként sorba rendezhetők. Konstans célváltozó alkalmazásával az objektum-azonosságok kényszerként jelennek meg a modellezési folyamatban, mely az anti-diszkriminációs számítások, vagyis a hasonlósági skála központi eleme (norma-értéke). A hasonlóságelemzés optimalizáló eljárás, amely minden egyes attribútumhoz egy lépcsős függvényt approximál, mely a célváltozóhoz történő hozzájárulás mértékét határozza meg (BÁNKUTI 2010). A lépcsők (idealitás-súlyszámok) közötti különbségek minimuma mindenkor nagyobb kell, hogy legyen, mint nulla, tehát különböző rangsorszámokhoz különböző lépcsős értékek kell, hogy tartozzanak. Hasonlóságelemzés által, így, a normától jelentősen eltérő objektumok ún. kiugró értékűként jelentkeznek, az objektumok célváltozó értéke szerint rangsorolhatóvá válnak. A hasonlóságelemzés anti-diszkriminatív eljárása egy speciális neurális hálónak, módszertanilag felügyelet nélküli gépi tanuló algoritmusnak tekinthető.

A hasonlóságelemzés anti-diszkriminatív eljárásának egyetlen „apró” változtatásával egy optimalizált termelési függvény-generáló eljárást kapunk. Amennyiben az adathalmaz tartalmaz metrikus célváltozót, a fiktív célváltozó helyettesítésével megfosztva a módszert annak anti-diszkriminatív jellegétől, egy felügyelt gépi tanuló algoritmus keretében, becsülhetővé válik a célváltozó, ahol a becslés és tény közötti eltérések értékelése egy kontextus független költség-haszon elemzésként értelmezhető (a lépcsős függvények felhasználásával). A termelési függvény súlyszámai között immár az egyenlőség is megengedett a hasonlóságelemzés anti-diszkriminatív verziójától eltérő módon. Minden anti-diszkriminatív hasonlóságelemzés képes tehát termelési függvényt generálni, de a lépcsők azonosságát megengedő termelési függvény-generáló hasonlóságelemzések nem képesek anti-diszkriminatív modellt építeni, hiszen kényszerűen értelmüket veszítik a lépcsős függvények ilyen esetben pl. a célváltozó-értékek összegét az objektumok és attribútumok szorzata által adott számmal elosztva, azaz átlagos hatásmértéket rendelve minden lépcsőszinthez. A hasonlóságelemzés képes a teljes szabályrendszer egyidejű optimalizálására.

A súlyszámok approximációjához definiálni szükséges a veszteségfüggvényt. Legyen x_i az adathalmaz i -edik objektuma, $S(x)$ a lépcsősfüggvény, mely bemenete az i -edik objektum, kimenete az objektumhoz tartozó súlyszámok. A veszteségfüggvény a tény és becslés négyzetes hibaösszegét minimalizálja, tehát:

$$L := \min \sum_{i=1}^n \left(\left(\sum_{j=1}^m S_j(x_i) \right) - Y_i \right)^2$$

A hasonlóságelemzés legfőbb előnyei:

- anti-diszkriminatív matematikára képes;
- alkalmas az inverz és direkt nézetek függvény-szimmetria ellenőrzésére, így az objektumok érvényesség-vizsgálatára;
- képes a teljes szabályrendszer egyidejű optimalizálására.

A hasonlóságelemzés hátrányai:

- mivel rangsorszámokkal operál, ezért van egy kényszerű információvesztés;
- a rangsorszámok feldolgozása miatt az objektumok és a rangsorszámok növekedésével a futásidő exponenciálisan nő.

A hasonlóságelemzés a genetikai potenciál-kereső algoritmusban és modell-preferencia levezetésben kerül felhasználásra.

3.2.7. Pearson-féle korreláció

A Pearson-féle korreláció (lineáris korreláció) metrikus attribútumok közötti lineáris kapcsolatok irányának meghatározására és összefüggésük mérésére alkalmas, ahol a korrelációs együttható (jele := r) fejezi ki a kapcsolat szorosságát a mért attribútumok között. A korrelációs együttható kiszámítása az alábbiak szerint történik:

$$r = \frac{\sum_{l=1}^n (x_l^i - \bar{x}^i)(x_l^j - \bar{x}^j)}{\sqrt{\sum_{l=1}^n (x_l^i - \bar{x}^i)^2 \sum_{l=1}^n (x_l^j - \bar{x}^j)^2}}$$

Az alkalmazott matematikai apparátusból következően, a korrelációs együttható értéke -1 és 1 között mozog. SAJTOS és MITEV (2007) alapján a korreláció erősségét az alábbi táblázat szolgáltatja (5. táblázat).

5. táblázat: A korrelációs együttható iránya és erőssége

Korrelációs együttható értékei	Kapcsolat iránya és erőssége
$r = 1$	Tökéletes pozitív kapcsolat
$0.7 \leq r < 1$	Erős pozitív kapcsolat
$0.2 \leq r < 0.7$	Közepes pozitív kapcsolat
$0 < r < 0.2$	Gyenge pozitív kapcsolat
$r = 0$	Nincs lineáris kapcsolat
$-0.2 < r < 0$	Gyenge negatív kapcsolat
$-0.7 < r \leq -0.2$	Közepes negatív kapcsolat
$-1 < r \leq -0.7$	Erős negatív kapcsolat
$r = -1$	Tökéletes negatív kapcsolat

Forrás: SAJTOS és MITEV (2007)

A Pearson-féle korreláció előnyei:

- metrikus változók közötti összefüggések erősségének és irányának meghatározására alkalmas;
- így pl. felhasználható a dolgozatban alkalmazott mátrixok attribútumainak irány-preferenciáinak visszaellenőrzésére.

A Pearson-féle korreláció hátrányai:

- kizárólag lineáris kapcsolat megállapítására alkalmas;
- nem alkalmas ok-okozati összefüggések feltárására.

Pearson-féle korreláció a modellezések során használt irány-preferenciák ellenőrzésében kerül felhasználásra.

3.2.8. Varianciaelemzés

A varianciaelemzés a számításban résztvevő attribútumok átlagai közötti eltérések mérésére és szignifikanciájára ad választ a varianciák vizsgálatával. Az F-próba alkalmazható az összefüggések meglétének igazolására, azaz a nullhipotézis tesztelésére, ahol a dolgozatban a társadalomtudományi kutatásokban általánosságban elvárt 0.05 valószínűségi értéket (szignifikanciaszintet) tekintem mérvadónak (UGRÓSDY 2018). A varianciaelemzés feltétele a függő változók normál eloszlása, valamint a szóráshomogenitás, mely utóbbit a Levene-teszt kiértékelésével vizsgálom. A Levene-teszt nullhipotézisének elvetése igazolja a szóráshomogenitás feltételének teljesülését. A dolgozatban egyszempontos varianciaelemzés alkalmazására van kizárólag szükség, mely azt jelenti, hogy az attribútumok faktora egytényezős.

A Varianciaelemzés előnyei:

- metrikus és nominális változók közötti összefüggések szignifikanciájának meghatározására alkalmas;
- az F-próba robosztus, egy-egy feltételezés nem teljesülése minimális hatással bír az első- és másodfajú hiba elkövetésének valószínűségére, ezért azok nem növelik meg számottevően az elkövethető hibás döntések volumenét (SAJTOS – MITEV 2007).

A Varianciaelemzés hátrányai:

- számos önkorlátozó feltételezéssel működik pl. szóráshomogenitás, normál eloszlás, homoszkedacitás, stb.;
- nem alkalmas ok-okozati összefüggések feltárására.

Varianciaelemzés a modellezések során használt egyes csoportok pl. modell típusok teljesítményének szignifikanciájának meghatározásában kerül felhasználásra.

3.3. Gépi tanuló rendszerek kiértékelése

3.3.1. Jóságmetrikák

A kutatás objektív eredményeinek megállapítására és a modellek jóságának mérésére az alábbi táblázatban (6. táblázat) ismertetett metrikákat alkalmazom, valamint a 3.3.2. és 3.3.3. alfejezetekben tárgyalt AUROC és AUPRC értékeket.

6. táblázat: Modellek értékelésére alkalmazott jóságmetrikák

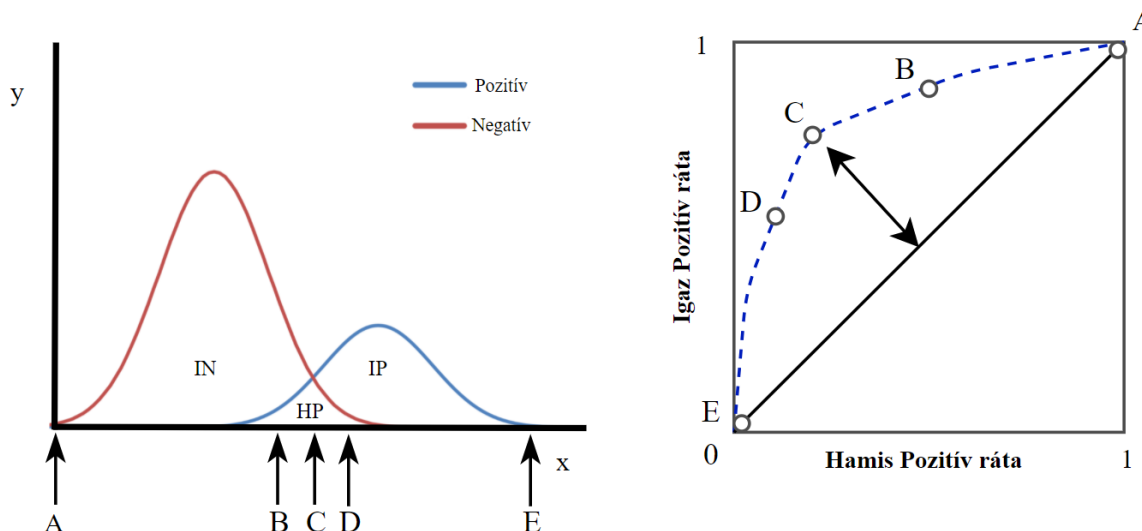
Jóság- metrika megnevezés	Leírás	Értelmezés
<i>Igaz pozitívák száma</i>	A modell igaz találatainak száma, amikor az helyesen eltalálta a fellépő gyanús kontrollt.	Értéke 0 és az összes kontrollhiányosság száma között mozog. Minél több az igaz pozitív találat, annál sikeresebb a hiányosságokkal rendelkező kontrollok felfedése. Irány-preferencia: minél nagyobb, annál jobb.
<i>Igaz negatívák száma</i>	A modell igaz találatainak száma, amikor helyesen eltalálta, hogy adott kontroll nem gyanús eset.	Értéke 0 és az összes megfelelő kontroll száma között mozog. Minél több az igaz negatív

		találat, annál sikeresebb a megfelelő kontrollok felfedése.
		Írány-preferencia: minél nagyobb, annál jobb.
<i>Hamis pozitívák száma</i>	A modell hamis találatának száma, amikor helytelenül arra a következtetésre jutott, hogy az adott kontroll gyanús eset, valójában pedig nem.	Értéke 0 és az adathalmaz elemeinek száma és igaz pozitív esetek különbsége között mozog. Minél kisebb a mutató, annál kevesebb a helytelenül gyanúsított kontrollok száma.
		Írány-preferencia: minél kisebb, annál jobb.
<i>Hamis negatívák száma</i>	A modell hamis találatának száma, amikor helytelenül arra a következtetésre jutott, hogy az adott kontroll nem gyanús eset, valójában pedig az volt.	Értéke 0 és az adathalmaz elemeinek száma és igaz negatív esetek különbsége között mozog. Minél kisebb a mutató, annál kevesebb a helytelenül nem gyanúsított kontrollok száma.
		Írány-preferencia: minél kisebb, annál jobb.
<i>Pontosság</i>	Az Igaz pozitívák és Igaz negatívák összegének és az összes eset hányadosa.	Értéke 0 és 1 között mozog. Minél nagyobb az értéke, annál nagyobb adott modell pontossága. A metrika kockázata, hogy kiegyensúlyozatlan osztályeloszlás esetén megtévesztő információval szolgáltat.
		Írány-preferencia: minél nagyobb, annál jobb.
<i>Precizitás (Precision)</i>	A Precizitás (Precision) az Igaz pozitívák számának, és az Igaz pozitívák és Hamis pozitívák összegének aránya. Kiegyensúlyozatlan adathalmazok esetén preferált mutató, melynek célkitűzése a modell pontosság becslésének relevanciájára vonatkozik a „Mennyire érvényes az eredmény?” kérdés megválaszolásával.	Értéke 0 és 1 között mozog. Minél nagyobb az értéke, annál nagyobb a helyesen gyanúsított kontrollok aránya az összes gyanúsított kontrollhoz képest.
		Írány-preferencia: minél nagyobb, annál jobb.
<i>Fedés (Recall)</i>	A Fedés (mely gyakran előfordul a magyar szakirodalomban még: felidézés, megbízhatóság és előhívás) az Igaz pozitívák számának, és az Igaz pozitívák és Hamis negatívák összegének aránya. Kiegyensúlyozatlan adathalmazok esetén preferált mutató, melynek célkitűzése a modell pontosság becslésének relevanciájára vonatkozik a „Mennyire teljes az eredmény?” kérdés megválaszolásával.	Értéke 0 és 1 között mozog. Minél nagyobb az értéke, annál nagyobb a helyesen gyanúsított kontrollok aránya az összes Igaz pozitív gyanúmomentumhoz képest.
		Írány-preferencia: minél nagyobb, annál jobb.
<i>F1-Pont</i>	Az F1-Pont a Precizitás és Fedés kombinációjából származtatott jószágmetrika, harmonikus középértéke:	Értéke 0 és 1 között mozog. Minél nagyobb az értéke, annál magasabb Precizitást (Precision) és Fedést (Recall) jelez.
		Írány-preferencia: minél nagyobb, annál jobb.
<i>Variancia</i>	A modell variancia mutatója fejezi ki a tanulás és tesztelési adathalmazon mért pontosságok közötti különbséget.	Értéke 0 és 1 között mozog. Minél kisebb az értéke, annál nagyobb a modell általánosító képessége, mivel előre nem definiált adatokon is képes közel hasonló performanciát nyújtani, mint a tanulóhalmazon.
		Írány-preferencia: minél kisebb, annál jobb.

Forrás: Saját szerkesztés

3.3.2. ROC-görbe és AUROC mutató

A bináris osztályozó algoritmusok performancia értékelésére széleskörben alkalmazott eljárás a ROC (Receiver Operating Characteristic - Vevő Működési Karakterisztika) görbék elemzése, valamint azok vizualizálása. A módszer alkalmazása lehetővé teszi az osztályozás igaz pozitív és hamis negatív találatok arányai közötti kompromisszum feltárását és értelmezését (TAN et al. 2018). A ROC görbék grafikus ábrázolása megköveteli az alkalmazott klasszifikáló eljárástól a becslések valószínűségének meghatározását, mely által az előrejelzések rangsorolhatóvá válnak a bizonyosság függvényében. Ezt szemlélteti a 14. ábra.



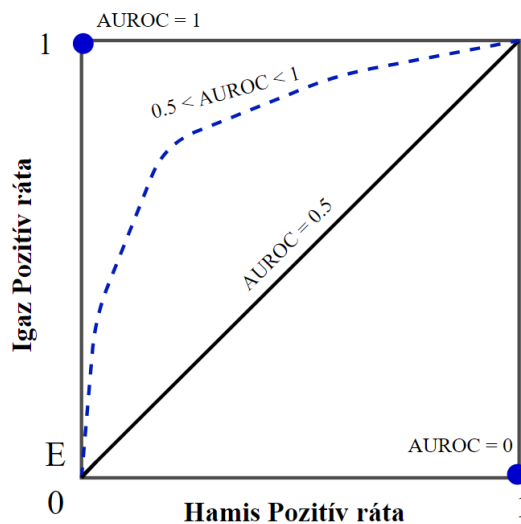
14. ábra: ROC-görbe vizualizálása

Forrás: Saját szerkesztés

Az AUROC (Area Under ROC – ROC Alatti Terület) mutatószám fejezi ki azt a ROC görbe határozott integrálásával meghatározott valószínűségi értéket, hogy egy véletlen minta pozitív osztályba tartozás becsült valószínűsége vélelmezhetően meghaladja-e egy véletlen minta negatív osztályba tartozásának valószínűségét. Minél jobban közelít az AUROC értéke 1-hez, annál alkalmasabb a pozitív és negatív osztályok különválasztására.

- AUROC = 1 esetén tökéletes osztályozásról beszélünk (azaz a ROC görbe hibátlanul illeszkedik a bal felső sarokba)
- AUROC = 0 esetén az osztályozó minden pozitív esetet negatívként, és minden negatív esetet pozitívként becsült.
- AUROC = 0.5 esetén véletlen becslésről beszélhetünk pl. ez akkor is előfordulhat, ha az algoritmus konstans módon minden mintát ugyanabba az osztályba sorol vagy irreleváns attribútum alapján történik az osztályozás, azaz az osztályozás szimplán a véletlen műve.
- $0.5 < \text{AUROC} < 1$ esetén a modell képes volt az adatban meghúzódó minták matematikai leképezésére, tehát több az igaz pozitív – igaz hamis találatok száma, mint a hamis pozitív és hamis negatív száma.

Az említett értékek szemléltetését vizualizálja a 15. ábra.

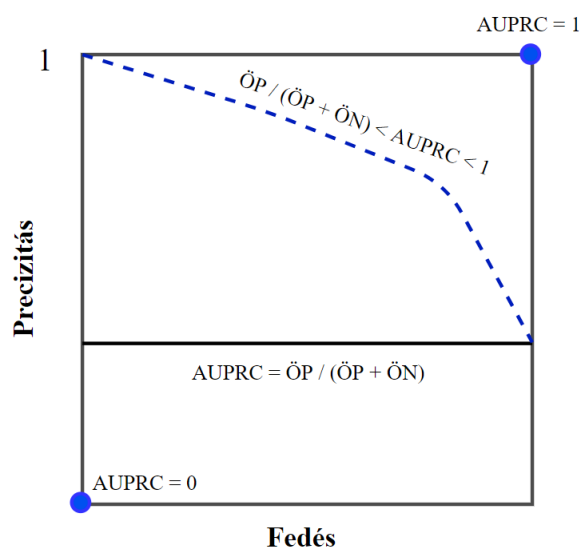


15. ábra: A ROC-görbe kitüntetett pontjai

Forrás: Saját szerkesztés

3.3.3. PR-görbe és AUPRC mutató

Kiegyensúlyozatlan osztályeloszlás esetén a ROC-görbék mellett alkalmazható eljárás a PR-görbék (Precision-Recall – Precizitás-Fedés) vizualizálása, mely a Precizitás és Fedés mutatók közötti kompromisszumot szemlélteti hasonlóan a ROC-görbékhez, mely esetén is értelmezhető a görbék alatti területek kiszámítása, az AUPRC (Area Under Precision-Recall Curve – PR-görbe Alatti Terület). Minél jobban közelít az AUPRC értéke 1-hez, annál magasabb adott modell Precizitás és Fedés értéke, mely rendre alacsonyabb hamis pozitív és hamis negatív arányt jelent. A 16. ábra szemlélteti a PR-görbét. Véletlen becslésről beszélhetünk, ha az AUPRC értéke megegyezik az összes pozitív (ÖP) és összes pozitív és összes negatív (ÖP + ÖN) minta hányadosával.

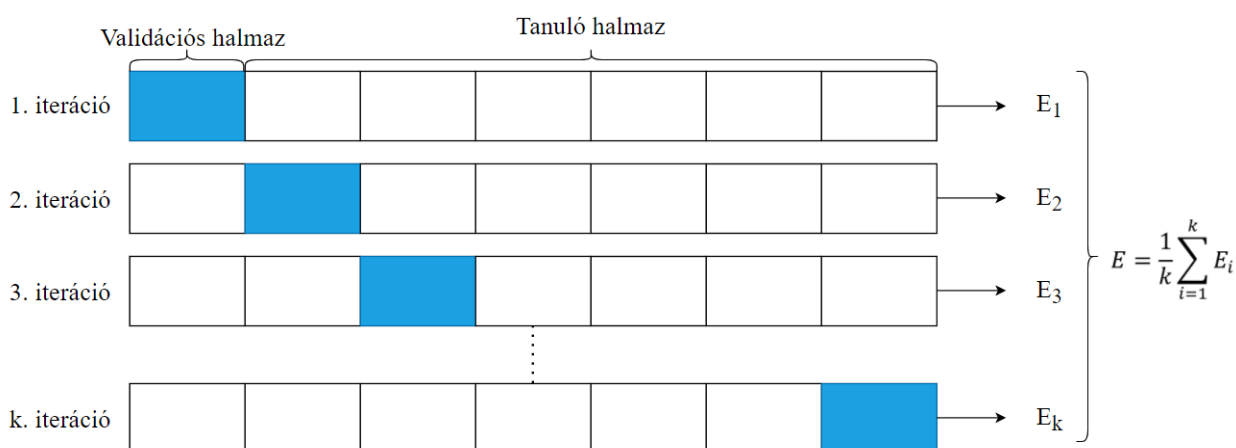


16. ábra: A PR-görbe vizualizálása

Forrás: Saját szerkesztés

3.3.4. Keresztvalidáció

Az alkalmazott modellek túltanulásának mérséklésére és az általánosító képesség megbízhatóbb érvényességének vizsgálatára a keresztvalidációs kiértékelés egy célszerű technikának bizonyul. A keresztvalidáció alkalmazásához a rendelkezésre álló adathalmazt k véletlen, egyenlőre-részre osztva minden minta legalább egyszer a validációs halmazba kerül, mely hozzájárul a modellek megfelelőbb variancia becsléséhez, valósabb képet ábrázolva a teljesítményekről. A végső teljesítménymutató meghatározása a k kiértékelések átlagaként értelmezhető, ahol az egyes iterációk részeredményei szolgáltatják az eljárás varianciáját. Minél nagyobb a keresztvalidáció varianciája, annál ingatagabb a modell, azaz a minta pl. annál zajosabbnak tekinthető. A k -iterációs (k -fold) keresztvalidáció alkalmazását vizuálisan a 17. ábra szemlélteti.



17. ábra: A keresztvalidáció működési logikája

Forrás: Saját szerkesztés

3.4. Felhasznált eszközök és technológiai megoldások

A dolgozatban bemutatott kutatás elkészítéséhez és modellezés megvalósításához a 7. táblázatban ismertetett eszközök kerültek felhasználásra.

7. táblázat: A kutatás során alkalmazott eszközök listája és leírása

Megnevezés	Kategória	Környezet/Verziószám	Leírás
Python	Programozási nyelv	Anaconda Navigator 1.9.12 Spyder fejlesztői környezet 4.0.1 Python 3.7.6	A Python egy célorientált programozási nyelv, mely támogatja a funkcionális és objektumorientált programozási paradigmákat. Az egyik legnépszerűbb nyíltforráskódú eszköz, mely kereskedelmileg is szabadon disztributálható, ezért egy kiterjedt fejlesztői közösséggel rendelkezik. Főbb felhasználási területei: hálózati biztonság (etikus hackelés) és adatelemzés.
Tensorflow	Szoftver könyvtár	2.3.0	A Tensorflow a Google által fejlesztett nyílt forráskódú szoftver könyvtár gépi tanuló alkalmazások programozására, melynek középpontjában a tenzorok és számítási gráfok állnak.

<i>Keras</i>	API	2.4.3	A Keras egy nyílt forráskódú Python API (Application Programming Interface – Alkalmazásprogramozási Interfész) a Tensorflow könyvtárhoz.
<i>Scikit-learn</i>	Szoftver könyvtár	0.22.1.	A Scikit-learn egy nyílt forráskódú szoftver könyvtár gépi tanuló rendszerek és statisztikai számítások programozására.
<i>MS Excel/VBA</i>	Táblázatkezelő és programozási nyelv	MS Office 2016 VBA 7.1.	Microsoft Office alkalmazásokban használt esemény-vezérelt programozási nyelv.
<i>MSSQL</i>	Adatbázis-kezelő	SQL Server 2016	Microsoft relációsadatbázis-kezelő rendszere, az SQL deklaratív nyelv segítségével hatékony adatelérést és adatmanipulációt tesz lehetővé.
<i>SPSS</i>	Szoftver	22.0.0.0	Az SPSS az IBM statisztikai szoftvercsomagja, mely alkalmazható egy- és többváltozós statisztikai számítások elvégzésére, alapvető gépi tanuló modellek elkészítésére.
<i>COCO</i>	Szoftver	v2.17	A hasonlóságelemzést futtató online elérhető eszköz megnevezése (component-based object comparison for objectivity - objektivitást támogató komponens-alapú objektum összehasonlítás)

Forrás: Saját szerkesztés

A dolgozattal kapcsolatos számításokat alátámasztó dokumentumok és programkódok elérhetők a Magyar Internetes Agrárinformatikai Újság serverén: <https://miau.my-x.hu/phdbg>, valamint a szerző GitHub Repository-jában: <https://github.com/AInside27/AI-methods-in-IT-security-audits>

3.5. A kutatási célok és hipotézisek rendszere

A 8. számú táblázat összefoglaló jelleggel ismerteti a kutatási célkitűzések és hipotézisek rendszerét, valamint az alkalmazott módszereket.

8. táblázat: Célkitűzések, hipotézisek és alkalmazott módszerek rendszere

Célkitűzések	Hipotézisek	Alkalmazott módszerek
C1: A Knuth-i elvet követve az információbiztonsági auditok hatékonyságát növelendő, létrehozandó olyan mesterséges intelligenciával ellátott döntéstámogató rendszer (robot-auditor), mely automatizáltan a historikus információbiztonsági auditjelentésekből tanulva képes a kontrollhiányosságok és kontrollterületek közötti összefüggések matematikai feltárására és javaslatteteleivel a potenciális emberi hibából fakadó észlelési kockázatok csökkentésére.	H1: Az információbiztonsági auditjelentések szöveges eredményeiből strukturált adatbázist alkotva és bemenetként a mesterséges intelligencia fogalmkörébe illeszthető eszközökkel azt feldolgozva, az auditok során feltárni kívánt kontrollhiányosságok megléte a véletlen találgatásnál nagyobb valószínűséggel kimutathatók, azaz a kontrollhiányosságok konstellációi matematikailag értelmezhető összefüggéseket hordoznak magukban. H1.1: A gyanúgenerálás, mint megoldandó üzletileg értelmezett probléma sajátosságait értékelve, a kontrollhiányosságok detektálása megoldható felügyelt és felügyelet nélküli gépi tanuló eljárásokkal is. H1.2: A gyanúgenerálás teljesítménye fokozható hibrid megközelítésben, azaz a felügyelt és nem felügyelt módszerek együttes felhasználásának a kutatásban alkalmazott releváns performancia metrikái ideálisabb értékeket mutatnak, mint önálló alkalmazásban H1.3: A hibrid modell többlet-információs értéket teremtve képes az egyszerű modellek általánosító képességén javítani.	Döntési fa Adaptív Boosting (ABM) Gradiens Boosting (GBM) Neurális háló Kollaboratív szűrés ROC analízis Varianciaelemzés
C2: A fejlesztendő mesterséges intelligenciával ellátott szoftveres robot-auditornak kényszerűen alkalmasnak kell lennie a rendelkezésre álló adathalmaz minél inkább az optimálishoz közeli felhasználására, mely által teljesítménye maximalizálható, azaz a cél a robot-auditor genetikai potenciáljának kiaknázása a tanulási adathalmaz irányított redukálása révén.	H2: A döntéstámogató rendszer genetikai potenciálja letapogatható hasonlóságelemzéssel ellátott kereső eljárással a tanításra alkalmazott adathalmaz irányított feldolgozásán keresztül, úgy, hogy a genetikai potenciálhoz vezető kereső eljárás a genetikus algoritmusok esetén alkalmazott véletlen mutáció és a populáció egyedeinek keresztezése nélkül is képes ideálisabb eredményt szolgáltatni.	Hasonlóságelemzés Pearson-féle korreláció Varianciaelemzés
C3: A tanulásra felhasznált adathalmaz információtartalmát növelendő, anti-diszkriminatív módon szükséges az egyes robot-auditor alternatívák teljesítményeinek összehasonlítása, a legjobb alternatíva kiválasztása, melyhez nem szükséges validációs és teszhalmaz elkülönítése a szokásos tesztelés általi adat/információ-vesztési gyakorlattal szemben.	H3: A mesterséges intelligenciával ellátott döntéstámogató rendszerek teljesítményalapon a gépi tanuló alkalmazások klasszikus tesztelési eljárásai nélkül is rangsorolhatók, a predikciók, mint generált gyanúforrások leíró tulajdonságainak érték-irány levezetésével és az ezen adatokat feldolgozó matematikai apparátussal, mely automatizáltan képes a preferált modellek objektív meghatározására.	Hasonlóságelemzés Pearson-féle korreláció Varianciaelemzés

Forrás: Saját szerkesztés

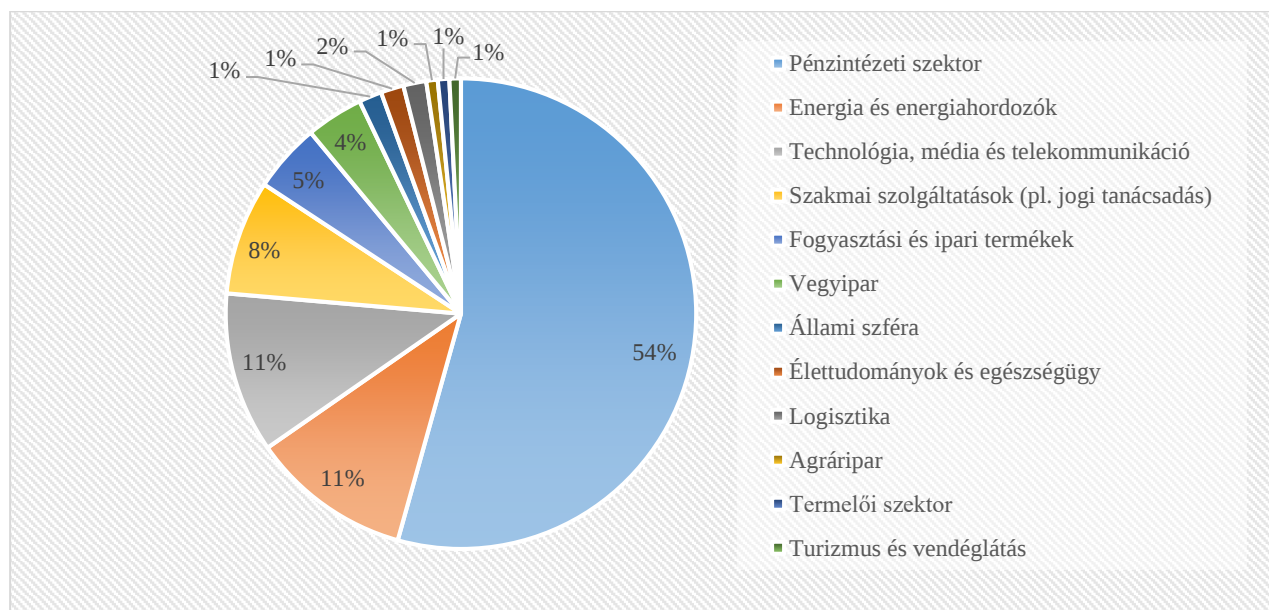
4. EREDMÉNYEK

A fejezet prezentálja a felállított hipotézisek bizonyítását a terepmunkán gyűjtött adatok elemzésére alapozva az előző fejezetben ismertetett matematikai apparátusok felhasználásával.

4.1. A kutatás során gyűjtött adatok leíró statisztikái

Az alfejezet a kutatás során gyűjtött adatvagyron leíró statisztikáit mutatja be, mely ezáltal magas szinten taglalja az auditok által vizsgált szervezetek, auditjelentések és kontrollmegfelelőségek karakterisztikáit.

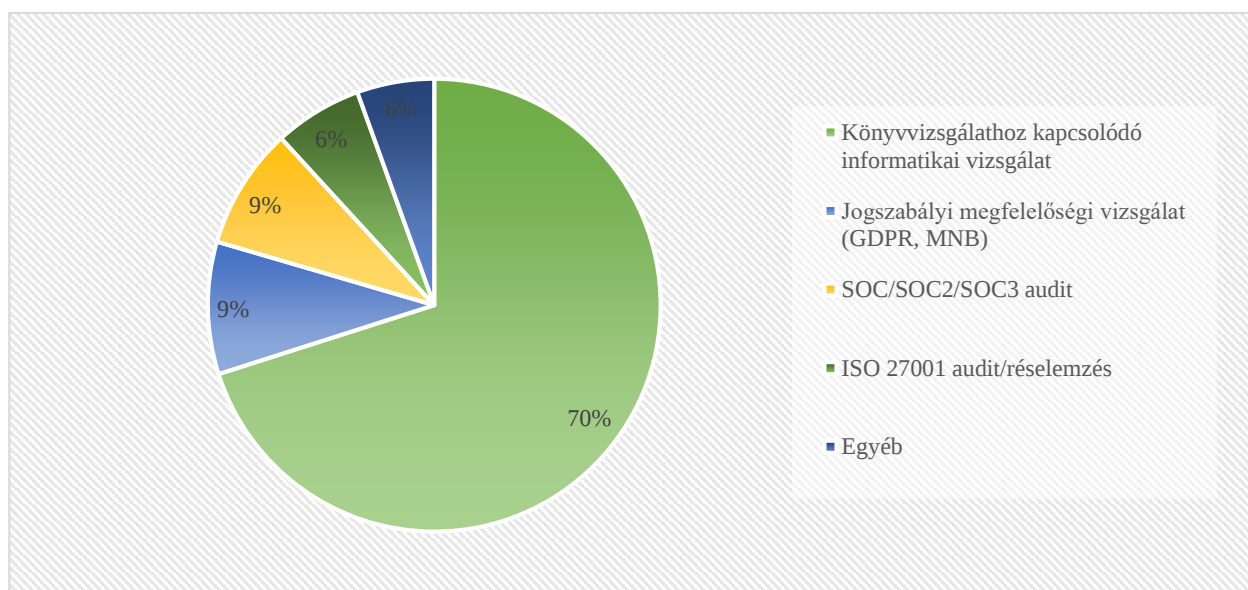
Az adatbázis 127 auditjelentés megállapításait tartalmazza iparáganként és audit típusonként, numerikus értékekkel kifejezve az audit megállapítások darabszámát az ISO/IEC 27001:2013 információbiztonsági szabvány „A” melléklete szerinti struktúrában. A 127 auditált szervezet megoszlását iparáganként a 18. ábra szemlélteti. Az ábrából jól kivehető, hogy túlnyomórészt a pénzügyi szektor (54%) állt a vizsgálatok hatókörében, mely köszönhető annak is, hogy az iparág a legjobban szabályozott iparágak egyike Magyarországon, és a Magyar Nemzeti Bank rendeleteiben és ajánlásaiban megköveteli a pénzügyi szervezetekre irányuló folyamatos független auditok elvégzését, például egy új banki technológia vagy termék bevezetésénél. Mivel az adatgyűjtés főként a pénzügyi szektorban történt, továbbá, a szervezetek méretével és árbevételével kapcsolatban nem állt rendelkezésre adat az anonimizáció miatt, ezért az alfejezetben ismertetett eredmények nem reprezentatívak, hiba lenne általánosítani a teljes Magyarországon működő vállalatok egészére, ezért szükségszerűen az egyes kiválasztott modellezési gyakorlatoknál (4.2.7. alfejezet) célszerű az adatbázis szűkítésével is elvégezni az elemzést. Fontos azonban kiemelni, hogy a kutatás során alkalmazott gépi tanuló algoritmusok robusztus eljárások, ezért, bár az adathalmaz heterogenitása kérdést vethet fel annak egyidejű felhasználásáról, az alkalmazott apparátusok gondoskodnak arról, hogy ez ne jelentsen problémát, szemben a többváltozó statisztikai eljárásokkal.



18. ábra: A gyűjtött minta megoszlása iparáganként

Forrás: Saját szerkesztés

A 127 gyűjtött audit típusainak megoszlását az alábbi ábra szolgáltatja (19. ábra). Az auditjelentések többségében könyvvizsgálathoz kapcsolódó informatikai vizsgálatok voltak (70%), melyet a jogszabályi megfeleléségi vizsgálatok (9%) és SOC auditok (9%) követtek.



19. ábra: A gyűjtött minta megoszlása audit típusonként

Forrás: Saját szerkesztés

Iparági megoszlásban a szervezetekre jutó megállapítások és hatókörben lévő kontrollok számát, továbbá a megállapítások és hatókörben lévő kontrollok arányát a 9. táblázat szemlélteti. A legtöbb szervezetre jutó megállapítás rendre a szakmai szolgáltatások (18.10 db), fogyasztási és ipari termékek (12.17 db), illetve a vegyiparban (10.40 db) volt tapasztalható. A legkevesebb kontrollhiányosság (1 db) a termelői szektorban került megállapításra, azonban, mivel kizárólag egyetlen auditjelentés állt rendelkezésre az iparágban, ezért az kiugró értéként kezelendő. A legtöbb megállapítással rendelkező szektor (szakmai szolgáltatások) kimagasló értéke a megállapítások vonatkozásában nem meglepő, mivel a legnagyobb hatókörben (105.90 db) történtek vizsgálatok, így értelemszerűen, nagyobb valószínűséggel kifogásolható az audit feltételezett kontrollmegfeleléseket.

A megállapítások és hatókörben lévő kontrollok arányának elemzése már beszédesebb, az oszlop az egy kontrollra jutó megállapítások számát fejezi ki, mely „győztese” a vegyipar (0.32 db), és így a szakmai szolgáltatások (0.17 db) a középmezőnyben szerepel. Az alacsony mintaszámmal rendelkező iparágakat leszámítva, a legkevesebb megállapítás arányosítva (0.10 db) a technológia, média és telekommunikáció szektorában jelenlévő szervezetekre volt jellemző, mely várható volt az információbiztonság fokozott követelményrendszere és az iparági sajátosságok együttes tükrében. Számomra meglepő módon a pénzügyi szektor magasabb megállapítási aránnyal rendelkezik (0.18 db), mint az előzetesen sejthető lett volna az erős jogszabályi megfeleléségi kényszer miatt, így is csak átlag (0.18 db) körüli eredményeket szolgáltat.

9. táblázat: Megállapítások, hatókörben lévő kontrollok száma, mutatói iparági megoszlásban

Iparágak	Megállapítások száma összesen (db)	1 szervezetre jutó megállapítások száma összesen (db)	Hatókörben lévő kontrollok száma összesen (db)	1 szervezetre jutó hatókörben lévő kontrollok száma összesen (db)	Megállapítások és hatókörben lévő kontrollok aránya (db/db)
<i>Agráripár</i>	2	2.00	30	30.00	0.07
<i>Állami szféra</i>	11	5.50	62	31.00	0.18
<i>Élettudományok és egészségügy</i>	11	5.50	60	30.00	0.18
<i>Energia és energiahordozók</i>	100	7.14	420	30.00	0.24
<i>Fogyasztási és ipari termékek</i>	73	12.17	249	41.50	0.29
<i>Logisztika</i>	14	7.00	60	30.00	0.23
<i>Pénzüntézet</i>	537	7.78	2910	42.17	0.18
<i>Szakmai szolgáltatások</i>	181	18.10	1059	105.90	0.17
<i>Technológia, média és telekommunikáció</i>	113	8.07	1077	75.93	0.10
<i>Termelői szektor</i>	1	1.00	30	30.00	0.03
<i>Turizmus és vendéglátás</i>	4	4.00	30	30.00	0.13
<i>Vegyipár</i>	52	10.40	165	33.00	0.32
Összesen (ill. átlagosan)	1099	8.65	6152	48.44	0.18

Forrás: Saját szerkesztés

Audit típusonkénti megoszlásban a szervezetekre jutó megállapítások és hatókörben lévő kontrollok számát, illetve a megállapítások és hatókörben lévő kontrollok arányát a 10. táblázat prezentálja. A legkiterjedtebb hatókörrel rendelkező vizsgálatok (114.00 db) az ISO/IEC 27001:2013 auditok és réselemzések voltak, melyek célja a standard megfelelőségének ellenőrzése, melyet a SOC auditok (89.55 db) és jogszabályi vizsgálatok (88.42 db) követtek. Legkisebb audit hatókörrel (30.33 db) a könyvvizsgálathoz kapcsolódó informatikai vizsgálatok rendelkeztek, mivel a könyvvizsgálati munkák kizárólag az információbiztonsági kontrollkörnyezet legmagasabb kockázatairól szándékoznak tudomást szerezni. A megállapítások arányait vizsgálva az audittípusok többé kevésbe az átlag körül (0.18) mozogtak, ettől kivétel a SOC vizsgálatok, ahol ez az arány a legkevesebb volt (0.10). A SOC auditok célközönsége főleg az auditált fél ügyfelei, ezért nem okozhat meglepetést az alacsony ráta, mivel az információbiztonsági kontrollok megfelelő operálása a legtöbb esetben szerződéses követelmény, így a szolgáltató szervezetek kényszerűen, a jó hírnév fenntartása és az ügyfelek bizalma érdekében magas színvonalon üzemeltetik a biztonsági környezetüket. Összességében elmondható a táblázat alapján, hogy minden 5. vizsgált kontroll esetén hiányosságot vélt felfedezni az audit.

10. táblázat: Megállapítások, hatókörben lévő kontrollok száma, mutatói audit típusonkénti megoszlásban

Audit típusok	Megállapítások száma összesen (db)	1 szervezetre jutó megállapítások száma összesen (db)	Hatókörben lévő kontrollok száma összesen (db)	Szervezetre jutó hatókörben lévő kontrollok száma összesen (db)	Megállapítások és hatókörben lévő kontrollok aránya (db/db)
Egyéb	88	12.57	495	70.71	0.18
ISO 27001 audit/réselemzés	147	18.38	912	114.00	0.16
Jogszabályi megfeleléségi vizsgálat	214	17.83	1061	88.42	0.20
Könyvvizsgálathoz kapcsolódó informatikai vizsgálat	556	6.25	2699	30.33	0.21
SOC/SOC2/SOC3 audit	94	8.55	985	89.55	0.10
Összesen (ill. átlagosan)	1099	8.65	6152	48.44	0.18

Forrás: Saját szerkesztés

Az 11. táblázat kontrollterületenkénti megbontásban részletezi a megállapítások és hatókörben lévő kontrollok számát és arányát, ahol az első oszlop a 2.2.5. alfejezetben ismertetett ISO/IEC 27001:2013 szabvány hivatkozási számait tartalmazza (a kontrollterületek rövid leírása a 2. számú mellékletben található). Legnagyobb hatókörben az A9. Hozzáférés szabályozás (1363 db) és az A12. A működtetés biztonsága (1158 db) kontrollterületek voltak, mely arra enged következtetni, hogy az adatbizalmasság követelménye és teljesülése az auditok egyik központi kérdése, mivel ezen kontrollterületek túlnyomórészt az adatbizalmasság kompromittálódásának kockázatának mérséklésére kívánnak kontrollokat definiálni. A legkedvezőtlenebb megállapítás aránnyal (0.71) az A5. Az információbiztonság vezetői irányítása kontrollterület rendelkezik, mely a szabályozási keretét határozza meg az információbiztonsági irányítási rendszereknek. Ez azt jelenti, hogy közel az esetek háromnegyedében az auditor hiányosságot állapított meg az információbiztonsági szabályzatok, előírások és azok felülvizsgálatának tekintetében, tehát a szervezetek alapvető dokumentációs és szervezési hiányosságokkal küszködnek, melynek nem triviálisan következménye a technikai kivitelezés kifogásolhatósága.

11. táblázat: Megállapítások, hatókörben lévő kontrollok száma, aránya kontrollterületenkénti megoszlásban

Kontrollterület azonosító	Hatókörben lévő kontrollok száma (db)	Megállapítások száma összesen (db)	Megállapítások és hatókörben lévő kontrollok aránya (db/db)
A5	73	52	0.71
A6	254	64	0.25
A7	128	30	0.23
A8	374	37	0.10
A9	1363	422	0.31
A10	48	13	0.27
A11	454	51	0.11
A12	1158	187	0.16
A13	386	33	0.09
A14	946	54	0.06
A15	266	26	0.10
A16	339	22	0.06
A17	156	73	0.47
A18	191	29	0.15
A19	16	6	0.38
Összesen (ill. átlagosan)	6152	1099	0.18

Forrás: Saját szerkesztés

A következő táblázat azt a 10 kontrollkövetelményt prezentálja, melyek minden egyes audit esetén a vizsgálat hatókörét képezték (12. táblázat), ahol a kontrollkövetelményben szereplő első számjegy a kontrollterület azonosítóját jelöli. A táblázatból leolvasható, hogy az A9. Hozzáférés szabályozás terület kontrollkövetelményei 9 esetben, bármilyen vizsgálatról is legyen szó, mindig az auditok tárgyát képezték. Ez jelentheti azt, hogy a szervezetek/auditorok a hozzáférés szabályozás kontrollterületén azonosítják a legtöbb hiányosságot, vagy a kontrollterületet tartják a legkockázatosabbnak, így az illetéktelen hozzáférések lehetőségét ellenőrizendő, kivétel nélkül szerepelt az auditokban.

A változáskezelés kontrollkövetelménye is megfigyelhető a táblázatban, mely az informatikai fejlesztések és módosítások (tehát változtatások) szabályozását célozza. Ennek oka eredhet az újfajta technológiai és rendszerimplementációs projektek megnövekedett számából, az agilis fejlesztési módszertanok népszerűségéből, valamint, kockázati oldalon értékelve, abból, hogy a szervezetek nem rendelkeznek hatékony változáskezelési folyamattal, mely az alábbi ellenőrzési pontokat foglalja magában:

- Változások üzleti igényének felülvizsgálata, jóváhagyása, költségelemzése;
- Változások fejlesztése és tesztelése;
- Változások üzleti hatásának elemzése;
- Változások független minőségbiztosítása;
- Változások éles üzemi-környezetbe történő implementálása és folyamatos nyomon követése.

12. táblázat: Top 10 kontrollkövetelmény az auditok hatókörében

Kontrollterület	Kontrollkövetelmény	Hatókör (auditjelentések száma db)
Hozzáférés szabályozás	9.1.2 Hozzáférés a hálózatokhoz és a hálózati szolgáltatásokhoz	127
Hozzáférés szabályozás	9.2.1 Felhasználók engedélyezése és törlése	127
Hozzáférés szabályozás	9.2.2 Felhasználói hozzáférés kiosztása	127
Hozzáférés szabályozás	9.2.3 A privilegizált hozzáférési jogok kezelése	127
Hozzáférés szabályozás	9.2.4 A felhasználók bizalmas hitelesítési információinak kezelése	127
Hozzáférés szabályozás	9.2.5 A felhasználói hozzáférési jogok felülvizsgálata	127
Hozzáférés szabályozás	9.2.6 A hozzáférési jogok eltávolítása vagy módosítása	127
Hozzáférés szabályozás	9.4.1 Információhoz való hozzáférés korlátozása	127
Hozzáférés szabályozás	9.4.2 Biztonságos bejelentkezési eljárások	127
A működtetés biztonsága	12.1.2 Változáskezelés	127
Összesen		1270

Forrás: Saját szerkesztés

A megállapítások száma szerint rangsorolva a kontrollterületeket (13. táblázat), 6 esetben megállapítható, hogy a legtöbb audit megállapítás az A9. Hozzáférés szabályozás területén történt, míg 2 esetben A12. A működtetés biztonsága, és 1-1 esetben A5. Az információbiztonság szervezete és az A6. Az információbiztonság irányítása területen volt.

A legtöbb megállapítás a privilegizált hozzáférési jogok kezelése esetén volt regisztrálva (116 db). A privilegizált jogok nem megfelelő kezelése véleményem szerint az egyik legnagyobb kockázati tényező, mivel a privilegizált felhasználók a legtöbb esetben képesek az informatikai erőforrások és adatok legnagyobb részéhez hozzáférni (adminisztrátori jogosultság), ezért a támadók elsődlegesen ezen felhasználói fiókok feltörését kísérlik meg, mert hozzáférést biztosíthat a teljes infrastruktúrához és adatvagyonhoz, ezért kiemelten fontos ezen felhasználók rendszeres ellenőrzése és kontrollálása.

A második legtöbb megállapítás a biztonságos bejelentkezési eljárások kontrollkövetelményt érinti (103 db). Amennyiben a kontroll sérül pl. gyenge jelszavak alkalmazása esetén, akkor az informatikai rendszerek hitelesített bejelentkezésként fogadhatnak egy illetéktelen személyt, mely következtésképp, engedélyezheti a támadó számára az informatikai hálózat belső feltárását.

A harmadik kritikus pont az eseménynaplózás kontrollkövetelménye (55 db), tehát az informatikai rendszerekben történt tevékenységek rögzítése. Naplózás hiányában a rendszerek feltörése után ellehetetlenedhet a visszaellenőrzés és az incidensek felderítése, ezért a három legtöbb megállapítással rendelkező kontrollterület együttes hiányossága kritikus biztonsági kockázatot jelent.

13. táblázat: Top 10 legmagasabb megállapítással rendelkező kontrollkövetelmények

Kontrollterület	Kontrollkövetelmény	Megállapítások száma (db)
Hozzáférés szabályozás	9.2.3 A privilegizált hozzáférési jogok kezelése	116
Hozzáférés szabályozás	9.4.2 Biztonságos bejelentkezési eljárások	103
A működtetés biztonsága	12.4.1 Eseménynaplózás	55
Az információbiztonság szervezete	6.1.2 Feladatok szétválasztása	53
Hozzáférés szabályozás	9.2.5 A felhasználói hozzáférési jogok felülvizsgálata	51
Az információbiztonság vezetői irányítása	5.1.1 Információbiztonsági szabályozás	44
A működtetés biztonsága	12.6.1 Technikai sebezhetőségek kezelése	43
Hozzáférés szabályozás	9.2.6 A hozzáférési jogok eltávolítása vagy módosítása	42
Hozzáférés szabályozás	9.2.2 Felhasználói hozzáférés kiosztása	39
Hozzáférés szabályozás	9.4.1 Információhoz való hozzáférés korlátozása	37
Összesen		583

Forrás: Saját szerkesztés

A következő táblázat ismerteti azon 10 kontrollkövetelmény rangsorát, ahol a legmagasabb volt az egy auditjelentésre eső megállapítások száma (14. táblázat).

Legelső helyen szerepel az A5. Információbiztonsági szabályozás (1.19), ahol átlagosan minden egyes olyan audit esetén, ahol a kontroll vizsgálva volt, az auditorok kifogásoltak legalább 1 nem megfelelőséget. A kontrollkövetelmény elvárása, hogy a szervezetek rendelkezzenek információbiztonsági szabályzatokkal, a szervezeti követelmények legyenek definiálva, a vezetés által jóváhagyva, szervezeten belül nyilvánosságra hozva és kommunikálva a munkavállalók és releváns külső partnerek felé. Amennyiben már szabályzati szinten sem megfelelő a követelmények és elvárások definiálása, megkérdőjelezhetővé válhat a technikai kikényszerítés és a gyakorlatok szabálykövető alkalmazása.

Továbbá, kitűnik, hogy bár a felhasználói jogosultságok felülvizsgálata esetén csak 51 db megállapítás született, ezek hatékonysága szintén kérdést vet fel szervezeti szinten, mivel a második helyen a privilegizált felhasználók hozzáféréseivel kapcsolatos hiányosságokat véltek az auditok felfedezni (0.91).

Kiemelendő, hogy a 4. és 5. rangsorszámmal rendelkező kontrollok működésfolytonossággal kapcsolatos hiányosságok, melyek az elmúlt év elején történt események, a COVID-19 miatti üzletmenetfolytonossági kiesések alátámaszthatnak. Megjegyzendő, hogy az auditjelentések is visszatükrözik a működésfolytonosságra vonatkozó ellenőrzések gyarapodását a 2020-as év második negyedévére.

14. táblázat: Top 10 egy auditjelentésre eső megállapítások száma kontrollterületenként

Kontrollterület	Kontrollkövetelmény	Hatókörben lévő kontrollok száma (db)	Megállapítások száma (db)	1 auditjelentésre eső megállapítások száma (db/db)
<i>Az információbiztonság vezetői irányítása</i>	5.1.1 Információbiztonsági szabályozás	37	44	1.19
<i>Hozzáférés szabályozás</i>	9.2.3 A privilegizált hozzáférési jogok kezelése	127	116	0.91
<i>Hozzáférés szabályozás</i>	9.4.2 Biztonságos bejelentkezési eljárások	127	103	0.81
<i>A működésfolytonosság információbiztonsági aspektusai</i>	17.1.3 Az információbiztonsági folytonosság ellenőrzése, felülvizsgálata és értékelése	41	31	0.76
<i>A működésfolytonosság információbiztonsági aspektusai</i>	17.1.1 Az információbiztonsági folytonosság tervezése	40	21	0.53
<i>Az információbiztonság szervezete</i>	6.1.2 Feladatok szétválasztása	116	53	0.46
<i>A működtetés biztonsága</i>	12.4.1 Eseménynaplózás	126	55	0.44
<i>Humán-erőforrás biztonsága</i>	7.2.2 Információbiztonsági tudatosság, oktatás és képzés	36	15	0.42
<i>Hozzáférés szabályozás</i>	9.2.5 A felhasználói hozzáférési jogok felülvizsgálata	127	51	0.40
<i>Megfelelőség</i>	18.2.1 Az információbiztonság független felülvizsgálata	25	10	0.40
Összesen (ill. átlagosan)		802	499	0.62

Forrás: Saját szerkesztés

Összesítve a kiértékelteket, a hipotézisek igazolásának szemszögéből jelentős többletinformáció az alábbi:

- Az adatvagyon legnagyobb részt a pénzügyi szektorra (54%) és a könyvvizsgálathoz kapcsolódó informatikai vizsgálatokra (70%) korlátozódik, ezért racionálisnak tekinthető az adatvagyon részhalmazainak elkülönített vizsgálata, hatásainak mérése a teljes adatvagyonra kivetítve (4.2.7. alfejezet);
- Legnagyobb mértékben az A9. Hozzáférés Szabályozás és A12. A működtetés biztonsága kontrollterületek álltak az auditok hatókörében, mely előbbi kategória esetén 3 kontroll is relativizálva magas megállapítás aránnyal rendelkezik (14. táblázat). Az A9. és A12. kontrollterületek, ezért, feltételezhetően nagyobb súllyal lesznek hatással a kontrollhiányosságok meghatározásában, melyet számításba vesz a 4.3. alfejezetben ismertetett kereső eljárás is.

4.2. Gyanúgenerálás információbiztonsági kontrollhiányosságok detektálására

Az alfejezet a terepmunka során gyűjtött adatokból történő tudás kinyerését, a legelső célkitűzés elérését, az alábbi hipotézisek igazolását prezentálja:

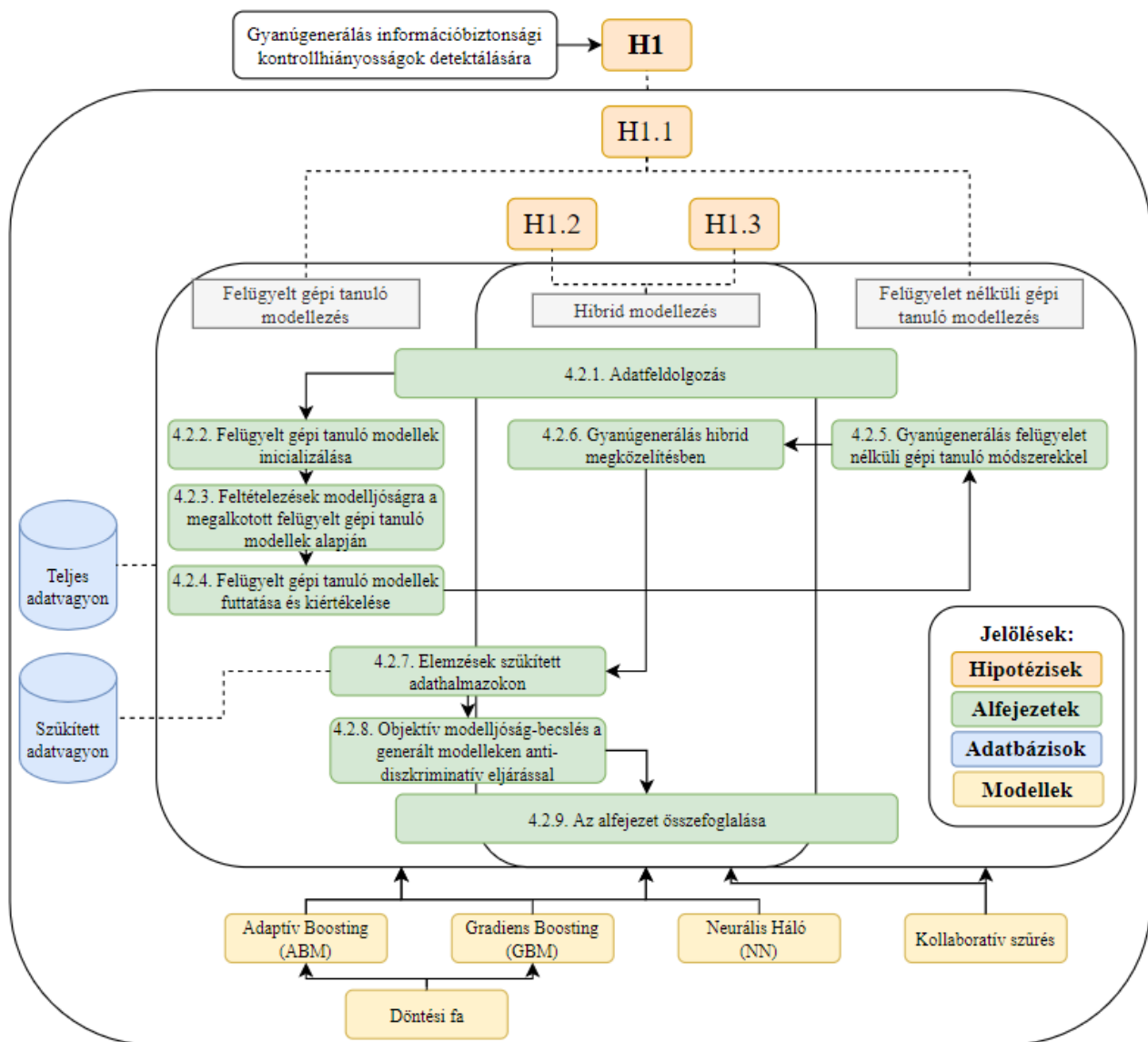
H1: Az információbiztonsági auditjelentések szöveges eredményeiből strukturált adatbázist alkotva és bemenetként a mesterséges intelligencia fogalmkörébe illeszthető eszközökkel azt feldolgozva, az auditok során feltárni kívánt kontrollhiányosságok megléte a véletlen találgatásnál nagyobb valószínűséggel kimutathatók, azaz a kontrollhiányosságok konstellációs matematikailag értelmezhető összefüggéseket hordoznak magukban.

H1.1: A gyanúgenerálás, mint megoldandó üzletileg értelmezett probléma sajátosságait értékelve, a kontrollhiányosságok detektálása megoldható felügyelt és felügyelet nélküli gépi tanuló eljárásokkal is.

H1.2: A gyanúgenerálás teljesítménye fokozható hibrid megközelítésben, azaz a felügyelt és nem felügyelt módszerek együttes felhasználásának a kutatásban alkalmazott releváns performancia metrikái ideálisabb értékeket mutatnak, mint önálló alkalmazásban

H1.3: A hibrid modell többlet-információs értéket teremtve képes az egyszerű modellek általánosító képességén javítani.

Az alfejezetben közöltek és hipotézisek kapcsolatát, valamint a hipotézisekhez tartozó modellezési gyakorlatokat az alábbi 20. ábra vizuálisan szemlélteti, rendszerezi. A kutatás nem kizárólag a teljes adatvagyon bevonásával történő gépi tanuló rendszerek fejlesztését és azok kiértékelését tűzte ki célul, hanem a szűkített, legnagyobb homogén csoportokat is külön kielemezte.



20. ábra: A 4.2. alfejezet rendszerezése

Forrás: Saját szerkesztés

4.2.1. Adatfeldolgozás

Elsődlegesen a terepmunka során gyűjtött adathalmaz strukturálására, valamint annak transzformációjára (dekódolás és standardizálás) volt szükség, hogy az adathalmaz, az alkalmazott algoritmusok számára feldolgozhatóvá váljon és így, alkalmas legyen az üzleti probléma megoldására.

A mindenkori cél egy adott kontrollt illetően előre jelezni, hogy a leíró ismérvek alapján minőségbiztosítási szempontból várható volt-e, hogy az audit folyamán hiányosságot vagy hatékonyságot fog az audit tapasztalni. Evégett, racionális az a döntés, hogy a modellezés célváltozója egy adott kontroll állapotát meghatározó attribútum legyen:

- Amennyiben, ha a célváltozó egy bináris változó, mely kifejezi, hogy a szóban forgó kontroll esetén létjogosultsága van-e egy adott megállapításnak, osztályozási problémáról beszélünk, ahol a negatív osztály a „nincs megállapítás”, tehát a kontroll hatékony, és a pozitív osztály a „van megállapítás”, azaz a kontroll nem hatékony.

- Lehetséges bináris változó helyett metrikus célváltozót is meghatározni (pl. hány darab megállapítás várható a kontrolltesztben), ebben az esetben pl. regressziós problémával van dolgunk.

A továbbiakban a gyanúgenerálást a kontrollhiányosságok detektálására vonatkozóan osztályozási problémaként azonosítom, így az alkalmazott algoritmusok azt az információt fogják szolgáltatni, hogy a tesztelendő kontroll a gyanú kategóriájába esik-e vagy sem. Mivel egy adott kontroll esetén üzleti szempontból a legjelentősebb eredmény az az, hogy adott kontroll hiányos-e vagy sem, mely hiányosságok számát egy regressziós algoritmus bár képes lehet tovább pontosítani, nagyobb az előrejelzés pontatlanságának a kockázata is, ezért költség/haszon megfontolásból elegendőnek ítélem a hiányosság meglétének előrejelzését a döntéselőkészítés során.

A célváltozóhoz tartozó tulajdonságok lehetnek a korábbi auditok folyamán feljegyzett, az adott auditra vonatkozó leíró karakterisztikák, mint pl. az auditban vizsgált szervezet iparági besorolása, az audit típusa, a hatókörben lévő rendszerek száma, és az auditban tapasztalt kontrollhiányosságok köre és száma. Mivel az audit iparági besorolása, típusa és a kontrollokra vonatkozó megállapítók ténye nominális változók, ezért igény mutatkozik azok transzformálására, melyhez a szakirodalomban is javasolt (pl. CERDA et al. 2018) OHE (One-hot encoding) módszer került kiválasztásra. Az adathalmaz összesen 12 különböző iparágban végzett audit vizsgálatot, 5 különböző audit típust és 115 kontrollt ír le a 3.1. és 4.1. alfejezetekben közöltekkel összhangban.

Az információtartalom teljes körű kihasználása végett egyéb mesterséges változók létrehozása is indokolt: így pl. a gyűjtött adatvagyonból származtatható az audit hatóköre (hatókörben lévő kontrollok száma), valamint a megállapítások és kontrollhatékonyságok száma. Ezen információt az ISO/IEC 27001:2013 standard alapján további 14 + 1 saját kontrollcsoportra lehet osztani. Következésképp, a célváltozóhoz tartozó attribútumok:

- 12 bináris változóban leírják, hogy milyen iparági auditról volt szó;
- 5 bináris változóban leírják, hogy milyen típusú auditról volt szó;
- 115 bináris változóban leírják, hogy mely kontroll esetén történik a célváltozó becslése;
- 45 bináris változóban leírják az auditra vonatkozóan annak hatókörét, megállapítások és megfelelések számát;
- továbbá, 1 diszkrét változóban a hatókörben lévő rendszerek számát.

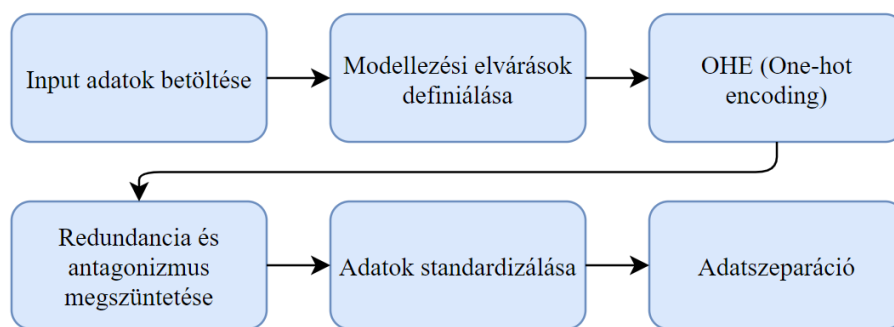
Mindezek fényében, összesen 178 különböző leíró attribútumot hoztam létre egy kontrollhiányosság megállapítására. Továbbá, a 127 feljegyzett auditból összesen 6152 olyan eset állapítható meg, ahol egy adott kontroll az audit hatókörét képezte, így a táblázat nyersen 6152 sort tartalmaz. Alapesetben 127x115, azaz 14605 sorról lenne szükség beszélni, azonban törlésre kerültek azon rekordok, melyekben az adott kontroll nem képezte az audit hatókörét, mivel nem várunk döntést a rendszertől olyan kontrollról, mely nem is volt vizsgálva, azaz, azzal a válasszal nem elégedhetünk meg, hogy az adott kontroll feltehetően nem hatókör, mert ez nem segíti az üzleti probléma megoldását, végső soron az audit munkáját.

A 6152 esetből adattisztítást követően 540 duplikátum volt tetten érhető, mely eredhet teljesen azonos auditok meglétéből, továbbá 5 esetben antagonizmus (ugyanazon attribútumok különböző célváltozó értékre vezettek) volt tapasztalható. 3 különböző esetben is, 6 duplikált értéknél az eltérő célváltozó értékek aránya 4:2-höz volt, ezért a kisebb számban lévő sorokat töröltem, míg két esetben 1:1-hez volt az arány, így mindkét rekord törlése került, mivel objektíven nem lehetett a rendelkezésre álló adatvagyon alapján eldönteni, hogy melyik rögzítés volt helyes vagy helytelen, vagy egyszerűen csak riportolási anomáliáról van szó. Az adattisztítást követően 532 duplikált értéket törölve, 5610 sor maradt a végleges adatbázisban, redundancia és antagonizmus

mentesen. A kollaboratív szűrésen alapuló felügyelet nélküli algoritmus esetén a bemeneti adattábla a 127 auditjelentést tartalmazta, így egyéb struktúra kialakítására nem volt szükség.

Mivel a neurális háló érzékeny az input adatok skálájára (pl. ennek függvényében lassabban konvergál), így az adatok standardizálása volt szükségszerű. A döntési fa alapú modellezés érzéketlen a skálázásra, ezért azon modellek esetében (ABM, GBM) a standardizálás minimális többlet hardverkapacitást igényelt, azonban többlethaszon (pl. jóságmetrikák javulása) nem keletkezett. Az alkalmazott kollaboratív szűrésen alapuló nem felügyelt módszer távolság/kapcsolat-metrikákon alapszik, így ezen alkalmazás tekintetében is indokolt volt a közös skálán történő operáció.

Az adattranszformációt (OHE és standardizálás) követően az adatokat tanulási és független tesztelési halmazra kettéosztottam (adatszeparáció) rendre 80% és 20% arányban. A tesztalmazban összesen 159 igazoltan kifogásolható, azaz hiányossággal rendelkező kontroll szerepel (14.17% - AUPRC küszöb), míg a hatékony kontrollok száma 963 (85.83%). Az algoritmusok technikai konfigurációjának letapogatása nem volt közvetlen célja a kutatásnak, így a keresztvalidációs halmazokon finomhangolást külön nem végeztem. Az adatfeldolgozás folyamatát az 21. ábra vizuálisan is szemlélteti.



21. ábra: Az adatfeldolgozás folyamata

Forrás: Saját szerkesztés

Indokolt megvizsgálni az adatszeparáció sikerességét, azaz a tesztalmaz megfelelően képes-e a teljes populációt reprezentálni. A 3. számú melléklet szemlélteti a teljes, tanuló-, valamint a tesztalmazra vonatkozó megoszlásokat audit típusonként és iparáganként. Megállapítható, hogy az elkülönített tesztalmaz megfelelően reprezentálja a rendelkezésre álló adatvagyon egészét.

4.2.2. Felügyelt gépi tanuló modellek inicializálása

A felügyelt gépi tanuló algoritmusok modellezési szakaszában a transzformált adathalmazt három különböző eljárással dolgoztam fel a 3.2. alfejezetben tárgyaltakkal összhangban:

- ABM (Adaptív Boosting) meta-tanulással ellátott döntési fa;
- GBM (Grádiens Boosting) meta-tanulással ellátott döntési fa;
- NN (Neurális Háló).

A kiválasztott módszerek alkalmazásának indoka a szakirodalmi áttekintő alapján levont konklúzió, miszerint az együttes módszerek (itt az ABM és GBM) meghaladják a naiv/alap/gyenge

algoritmusok teljesítményét (pl. jóságmetrikáit), valamint a neurális háló univerzális approximátor, ezért azaz elvárás, hogy magas színvonalon alkalmazható a probléma megoldására.

Az alkalmazott technikai konfigurációt a 15. táblázat szemlélteti, melyet MEASE és WYNER (2008), HASTIE et al. (2017), RASCHKA és MIRJALILI (2019) és GOODFELLOW et al. (2016) ajánlásaival összhangban határoztam meg.

15. táblázat: Felügyelt gépi tanuló eljárások technikai konfigurációi

Eljárás	Konfiguráció
<i>Döntési fa</i>	Kritérium: entrópia (legjobb hasítás) Maximális mélység: 3 Minimális mintavágás: 2 Minimális minta egy levélben: 1 Mintasúlyozás: nincs Osztálysúlyozás: nincs
<i>Neurális háló</i>	Rejtett rétegek száma: 2 Rejtett neuronok száma: 120 Tanulási ráta: 0.01 Tanulási optimalizációs eljárás: adaptív momentum Aktivációs függvény: relu (rectified linear unit – módosított lineáris egység) Kötegek feldolgozása: 200 Maximális iteráció: 200
<i>Adaptív Boosting</i>	Iterációk száma: 250 Tanulási ráta: 0.1
<i>Gradiens Boosting</i>	Iterációk száma: 250 Veszteségfüggvény optimalizáló eljárás: deviancia Tanulási ráta: 0.07

Forrás: Saját szerkesztés

4.2.3. Feltételezések modelljóságra a megalkotott felügyelt gépi tanuló modellek alapján

Az ABM algoritmus alapvetően érzékeny a kiugró értékekre, mivel a kivételeket súlyozva hoz ismételt döntést iterációról iterációra (megváltoztatva a minta eloszlását), ezért bár magas igaz pozitív találattal rendelkezhet, kiegyensúlyozatlan osztályozási probléma esetén (amilyen a rendelkezésre álló adathalmaz is) ez magas hamis pozitív aránnyal párosulhat, így hajlamos a túlilleszkedésre. A GBM nagy előnye annak flexibilitása, mely a pseudo-reziduálisok tanulásából ered, ezért többféle probléma esetén alkalmazható az ABM-mel ellentétben, mely bináris klasszifikációs problémák megoldására lett kifejlesztve, tehát a kutatás középpontjában álló problémát illene mindkét algoritmusnak elfogadhatóan abszolválnia.

Mivel a GBM a szekvenciális feldolgozás során az összesített hiba minimalizálására törekszik, feltételezhető, hogy a megoldandó probléma esetén magas igaz negatív találattal fog rendelkezni (közel 86%-a a teszhalmaznak a negatív osztályba tartozik). A Boosting algoritmusok által alkalmazott alaposztályozó merőben meghatározza azok teljesítményét, mely a kutatás esetén egy 3 mélységű döntési fa, s így, mivel az alaposztályozó kevésbé komplex struktúrát alkalmaz,

véleményem szerint mérsékelhető a modellek varianciái. A neurális háló legnagyobb gyengesége a konfigurálható paraméterek száma és az algoritmus adatéhsége, így az NN esetében magas varianciára számítok. Még az adatok és a paraméterek ismeretében is nehezen állítható fel vélhető rangsor a teljesítményre vonatkozóan, mivel mind a három modell használhatósága empirikus kísérletezéshez köthető (minden algoritmushoz tartozik olyan adathalmaz, melyet az old meg a legjobban), valamint mind a három módszer gyengesége a fekete doboz jellege, tehát a döntési logika kevésbé lesz számon kérhető, mely lényegében nem volt célja a dolgozatban bemutatott kutatásnak.

4.2.4. Felügyelt gépi tanuló modellek futtatása és kiértékelése

Az algoritmusok fejlesztését és „üzembehelyezését” követően a 16. táblázatban ismertetett, a független tesztelési halmazon mért eredményeket rögzítettem, melyet a keresztvalidációs kiértékelés szolgáltatott. A legtöbb igaz pozitív találatot az ABM produkálta (80 db), majd az NN (74 db), legvégül a GBM, mely csak 46 esetben volt képes igaz pozitív találatot teljesíteni, ezzel nagyságrenddel a két másik alkalmazott módszer performanciája alatt maradt, mely a többi mutatóban is tetten érhető. Azonban, az igaz negatívok számát vizsgálva a GBM (954 db) felülkerekedett az ABM (919 db) és NN (916 db) algoritmusokon, mely vélelmezhetően azt támasztja alá, hogy a GBM alapértelmezetten a mintákat (kontrollokat) megfelelőnek értékelte, alacsony hamis pozitív találatokat produkált. Az F1-Pont fejezi ki a Fedés és Precizitás együttes harmóniáját egy 0-tól 1-ig terjedő skálán (melyet 100-zal megszorozva százalékos teljesítményt kapunk, hasonlóan a Pontosság, Precizitás, Fedés, AUROC és AUPRC esetén), mely az ABM (0.57) alkalmazásnál a legkedvezőbb. Az ABM és GBM AUROC mutatói (0.85) a legideálisabbak, míg az AUPRC-t vizsgálva ez a GBM esetén állapítható meg (0.61), valamint a GBM modell rendelkezik a legkisebb varianciával (0.04).

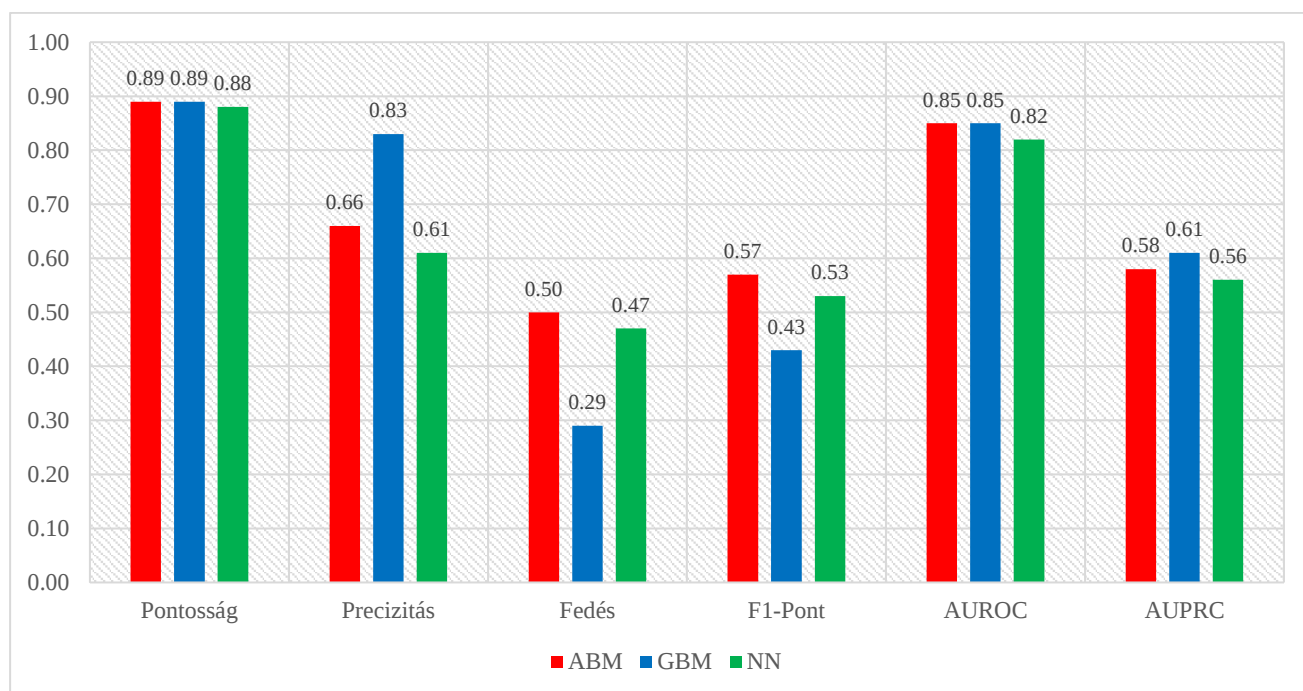
Összességében, az F1-Pont mutatót vizsgálva az ABM eljárás nevezhető ki a három algoritmus közül a legkielégítőbbnek, azonban ez vezetői döntéshez kötött. Amennyiben a cél a rendelkezésre álló erőforrások függvényében az igaz pozitív találatok maximalizálása annak árán is, hogy a rendszer számos hamis pozitívat generál, az ABM implementálása indokolt.

16. táblázat: Felügyelt gépi tanuló algoritmusok performancia metrikái

Performancia mutatók	Adaptív Boosting (ABM)	Gradiens Boosting (GBM)	Neurális Háló (NN)
<i>Igaz pozitívok száma (db)</i>	80	46	74
<i>Igaz negatívok száma (db)</i>	921	954	916
<i>Hamis pozitívok száma (db)</i>	42	9	47
<i>Hamis negatívok száma (db)</i>	79	113	85
<i>Pontosság</i>	0.89	0.89	0.88
<i>Precizitás</i>	0.66	0.83	0.61
<i>Fedés</i>	0.50	0.29	0.47
<i>F1-Pont</i>	0.57	0.43	0.53
<i>Variancia</i>	0.08	0.04	0.11
<i>AUROC</i>	0.85	0.85	0.82
<i>AUPRC</i>	0.58	0.61	0.56

Forrás: Saját szerkesztés

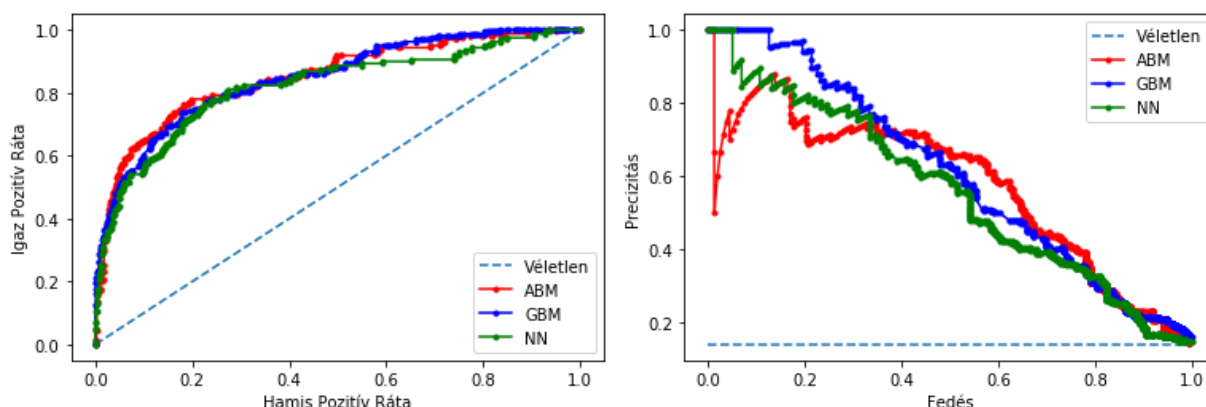
A modellek kiválasztott teljesítménymutatóit az 22. ábra vizuálisan prezentálja, ahol az Y tengely 0-tól 1-ig terjedő skálán szemlélteti a Pontosság, Precizitás, Fedés, F1-Pont, AUROC és AUPRC metrikákat.



22. ábra: Felügyelt gépi tanuló algoritmusok performancia metrikái

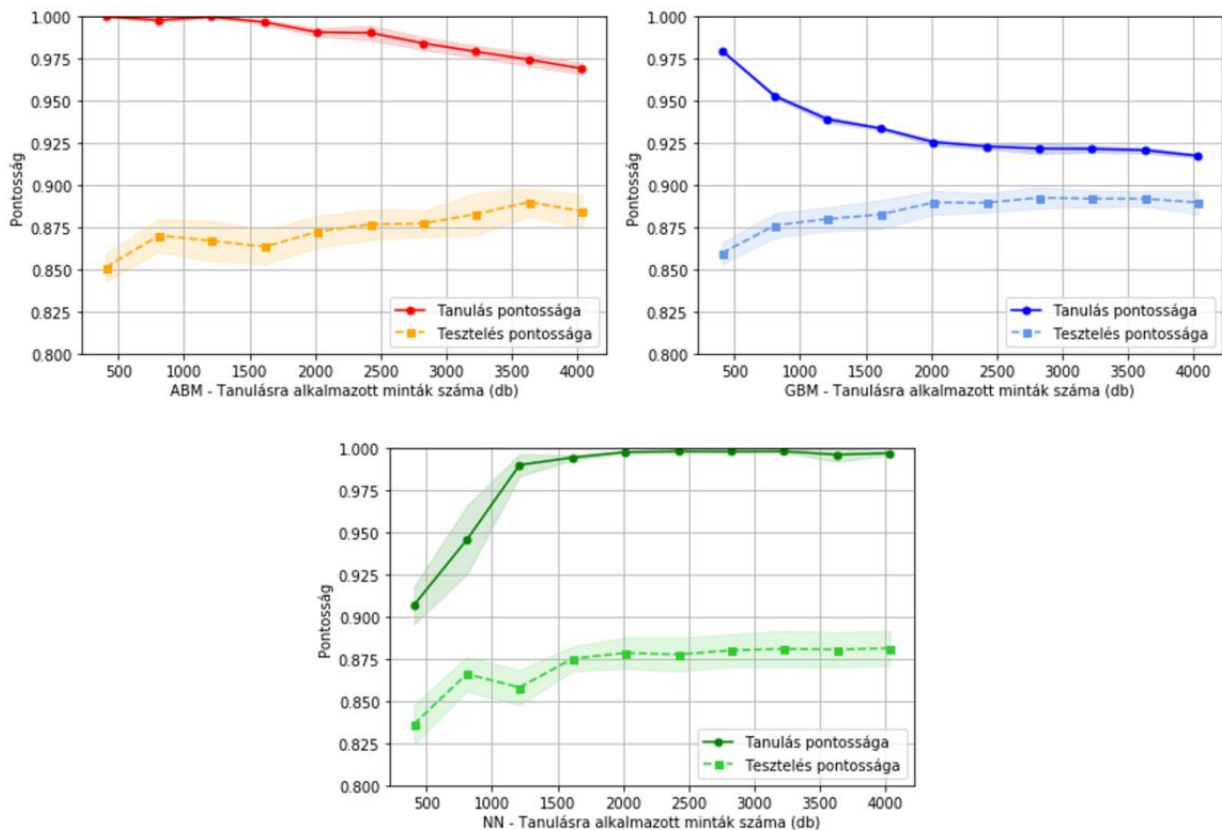
Forrás: Saját szerkesztés

A ROC-görbék és PR-görbék grafikus ábrázolásait a 23. ábra szolgáltatja. A görbéken is szemrevételezhető módon elhanyagolható különbség látható a modellek által lefedett területek nagysága között, nem tapasztalható jelentősen kiugró görbe, tehát az igaz pozitív és hamis pozitív találatok, valamint a Precizitás és Fedés közötti kompromisszum hasonlóan alakul valamennyi modell esetén.



23. ábra: Felügyelt gépi tanuló modellek ROC és PR-görbéi

Forrás: Saját szerkesztés



24. ábra: Felügyelt gépi tanuló algoritmusok tanulási görbéi

Forrás: Saját szerkesztés

A 24. ábra az alkalmazott algoritmusok tanulási görbéit szemlélteti a Pontosság tükrében, melyet fenntartásokkal kell kezelni a kiegyensúlyozatlan osztályeloszlás miatt, mindazonáltal hasznos információt szolgáltatnak a túlilleszkedés vizsgálatára. Az X tengely a tanulásra felhasznált minták darabszámát, az Y tengely a Pontosságot ábrázolja. Az egyes vonaldiagrammok a keresztvalidáció középértékeit, az azokat körbe ölelő sávok a becslések variánsait mutatják.

Általánosságban megállapítható, hogy mind a három módszer túlilleszkedett, azaz az általánosító képességük nem optimális, a tanulási halmazon mért pontosság meghaladja a tesztelési metrika eredményét, tehát az algoritmusok az adatban meghúzódó zajokat, valamint egyéb a kizárólag a tanulási halmazra jellemző karakterisztikákat is beépítettek az approximációba, mely várható volt. A variancia értéke a legkisebb a GBM esetében (0.04), míg a legnagyobb az NN modellben (0.11), ahol a tanulás 2000 db minta környékén szinte maximális (az algoritmus konvergált), tehát a modell a tanulóhalmazon minden mintát helyesen értékelt, ami a magas túlilleszkedés és alacsony általánosító képesség jeleit vélelmezik, ezért éles környezetben megkérdőjelezhető az alkalmazás fenntarthatósága. A túlilleszkedés mértéke csökkenthető addicionális tanulási adatok beszerzésével, valamint a modellek regularizációjával.

Összefoglalva a kiértékelteket, az AUROC és AUPRC eredményeit vizsgálva kijelenthető, hogy mind a három alkalmazott módszer képes volt a kontrollhiányosságok között meghúzódó mintázat matematikai leképezésére, mivel a kívánt véletlen szint felett teljesítettek (az AUROC értéke meghaladta a 0.5-et, az AUPRC értéke a 0.14-ot). Továbbá, mivel az F1-Pont értéke értelmezhető és értéke nagyobb nullánál (az algoritmusok nem minden rekordot egy megadott osztályba

soroltak), így kijelenthető, hogy a rendszerek teljesítménye jobb a véletlen találgatásnál, a kontrollhiányosságok együttes megléte között összefüggés állapítható meg.

4.2.5. Gyanúgenerálás felügyelet nélküli gépi tanuló módszerekkel

A felügyelet nélküli módszerek is alkalmazhatók gyanúgenerálása, azonban a teljesítmények mérése körülményes, mivel nincs visszaigazolt célváltozó az eljárások sajátosságainak köszönhetően. A kollaboratív szűrés és a gyanúgenerálás között párhuzam vonható, tehát a gyanúgenerálásra alkalmazott döntéselőkészítő rendszer, ajánlórendszerként is értelmezhető, ahol az auditjelentés számára kell mesterségesen a historikus adatok alapján gyanús kontrollt „ajánlani”.

A gyanúgenerálás felügyelet nélküli esete a hasonlóságelemzés célváltozó nélküli rétegeként is értelmezhető, ahol a gyanú matematikai jelenség: az irányított input rétegek kapcsán, ha nem igaz, hogy minden objektum lehet „másképp egyforma”, akkor a nem normaszzerű objektumok gyanúsak.

A kísérletben a 3.2.5. alfejezetben bemutatott algoritmust alkalmaztam, ahol az algoritmusok esetén három különböző távolság/kapcsolat-metrikát használtam. Jelölje d a távolság/kapcsolat-függvényt és x a kiválasztott objektumokat:

$$(1) \text{ EUC} = \text{Euklideszi távolság: } d(x^i, x^j) = \sqrt{\sum_{l=1}^n (x_l^i - x_l^j)^2}$$

$$(2) \text{ PEA} = \text{Pearson-féle korreláció: } d(x^i, x^j) = \frac{\sum_{l=1}^n (x_l^i - \bar{x}^i)(x_l^j - \bar{x}^j)}{\sqrt{\sum_{l=1}^n (x_l^i - \bar{x}^i)^2 \sum_{l=1}^n (x_l^j - \bar{x}^j)^2}}$$

$$(3) \text{ COS} = \text{Koszinusz hasonlóság: } d(x^i, x^j) = \frac{\sum_{l=1}^n x_l^i x_l^j}{\sqrt{\sum_{l=1}^n x_l^i{}^2} \sqrt{\sum_{l=1}^n x_l^j{}^2}}$$

Az alkalmazott eljárások során a döntéshozó preferenciája az, hogy nyilatkozzon az elfogadható hasonlóság küszöbértékéről, vagy azon minimum és maximum auditjelentések számáról, melyek a legközelebb találhatók a kiválasztott objektumhoz a többdimenziós térben. A kutatás alkalmával önkényes módon az utolsó lehetőséget választottam, ahol minden tesztelésre kiválasztott auditjelentés esetén 10-ben maximalizáltam a leghasonlóbb objektumok számát, így minden egyes futási eredmény egy meghatározott auditjelentés tekintetében egy 9 elemű listával tért vissza (plusz az objektum önmaga, mely értelemszerűen üres eredményt, azaz nulla különbséget szolgáltat).

A visszatérési értékek közötti redundanciát megszüntettem egy közös listára történő aggregációval, tehát ezen lista határozza meg a rendszer által gyanúsak vélt kontrollokat (igaz osztályozás). A teljesítmény mérhetőségének feltételeként az algoritmusokat egy szimulált környezetben alkalmaztam, ahol véletlenszerűen egy pozitív létező kontroll értékét megváltoztattam. Amennyiben az algoritmusok képesek a mesterségesen módosított rekord megtalálására, azt pozitív találatként könyveltem el, minden más gyanúmomentumot hamis pozitívként – noha elvileg a leghasonlóbb objektumok maguk is lehetnek abban az értelemben hibásak, hogy esetükben nem minden megállapítás született meg, aminek illett volna. A lista komplementer halmaza adta így az igaz negatívakat, míg, ha az algoritmus nem találta meg a

módosított értéket, azt hamis negatívként rögzítettem. Bár a teljesítmény efféle mérése nem biztosít teljes mértékben megbízható performancia adatokat, közelítő megoldásként elfogadható, azzal az előzetes feltételezéssel, hogy várhatóan magas lesz a hamis pozitív találatok száma szisztematikusan.

A Koszinusz hasonlósággal kalkuláló algoritmus két vektor által bezárt szög alapján hoz ítéletet, tehát ugyanazon hatókörrel rendelkező auditjelentések között azon jelentések lesznek a leghasonlóbbak, ahol megegyező kontrollok között volt tapasztalható hiányosság, a hiányosságok számától kvázi függetlenül, mely az Euklideszi távolsággal operáló módszertől eltérő, mivel két távoli kontroll alapján is lehet egy jelentés közeli, ha pl. alacsony a megállapítások száma. A Pearson-féle korreláció két rekord közötti lineáris kapcsolat erősségét vizsgálja, és megegyezik a Koszinusz hasonlósággal, amennyiben az objektumok átlagai nulla (tehát ez adja a kettő közötti különbséget), mely alapján, vélhetően, a PEA és COS közel hasonló eredményt fog szolgáltatni, és feltételezhetően jobbat, mint az EUC (ha minimális azon auditjelentések száma, melyben nem volt tapasztalható egyáltalán megállapítás, mely PEA és COS esetén nullát eredményez a nevezőben). Vélhetően, az eltérés az algoritmusok igaz pozitív találatok között minimális lesz az alacsony mesterségesen létrehozott pozitív osztály elemszáma miatt.

Tesztelésre 25 véletlen minta került kiválasztásra, mely a teljes populáció (127) 19.69%-a. Az algoritmusok futtatását követően a 17. táblázatban ismertetett eredményeket rögzítettem.

17. táblázat: Felügyelet nélküli gépi tanuló algoritmusok performancia metrikái

Performancia mutatók	EUC	PEA	COS
<i>Igaz pozitívak száma (db)</i>	16	18	20
<i>Igaz negatívak száma (db)</i>	2739	2721	2710
<i>Hamis pozitívak száma (db)</i>	111	129	140
<i>Hamis negatívak száma (db)</i>	9	7	5
<i>Pontosság</i>	0.96	0.95	0.95
<i>Precizitás</i>	0.13	0.12	0.13
<i>Fedés</i>	0.64	0.72	0.80
<i>F1-Pont</i>	0.21	0.21	0.22

Forrás: Saját szerkesztés

A kiválasztott minta esetén az alkalmazott módszertan függvényében, az algoritmusok összesen potenciálisan 25 gyanús kontrollról voltak képesek a feltételezett kontrollhiányosságokról igaz pozitív véleményt alkotni. Ebből a COS (20 db) 80%-át, a PEA (18 db) 72%-át és az EUC (16 db) 64%-át tudta helyesen megítélni a könnyen előre megjósolható magas hamis pozitív találati arány mellett. A Fedés metrika, mely az igaz és hamis pozitívak arányát fejezi ki, kedvező értéket mutat mindhárom modell esetében, ami a mesterségesen kikényszerített látens célváltozó alkalmazásának köszönhető. Mivel a visszatérési listák az összes eltérést tartalmazzák mely a kiválasztott objektumok között felmerül, ezért a Precizitás alacsony szintje nem meglepő. Az F1-Pont közel azonos értéke nem enged szignifikáns különbségre következtetni az eljárások között, minimális differenciával a COS tekinthető az F1-Pont aspektusából a legjobb módszernek.

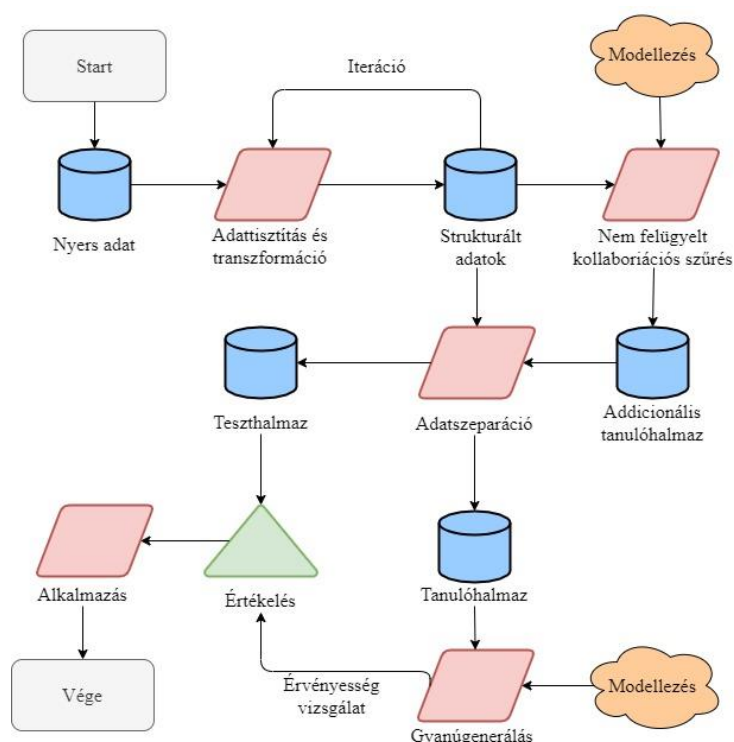
Kiemelendő, hogy a mérés szimulált környezetben történt, ezért a hasonló objektumok a felállított virtuális valóság nélkül akár kevésbé tűnhetnek hasonlóknak, sőt, az is elképzelhető, hogy egyáltalán nem szerepelt volna egy kitüntetett kontroll a kiválasztott tesztelési minta hiányosságokat tartalmazó listájában. Továbbá, fizikai mérőeszköz hiányában az eredeti adatok

rögzítésének pontossága is potenciális kockázat, így annak is van létjogosultsága, hogy az algoritmusok olyan kontrollhiányosságokat észleltek, melyeket az auditorok korábban hibásan megfelelőnek minősítettek, azonban ennek visszaellenőrzésére sajnos nincs mód. Az alkalmazott felügyelet nélküli modellek alkalmazása így egyelőre kockázatos a visszacsatolás hiánya miatt (mint általában bármely klaszterező eljárás eredményének kiértékelése), így a módszer hozzáadott értékéről a hibrid feldolgozás során nyerhetünk további tanúbizonyságot (következő alfejezet).

4.2.6. Gyanúgenerálás hibrid megközelítésben

A hibrid modellezés alapmegközelítése a releváns szakirodalmat felhasználva és következtetéseket levonva, hogy az előző alfejezetben ismertetett felügyelet nélküli algoritmusok segítségével növeljük a becslések jószágmetrikáit, tehát a felügyelt módszerekbe beépítve az „ajánlórendszer ajánlásait”, azok további hasznos bemenetre tegyenek szert, vélelmezhetően addicionális összefüggéseket tartalmazva.

A hibrid modellezésben a COS modell kimeneti értékeit építettem be a felügyelt módszerekbe (ABM, GBM, NN) ceteris paribus, a felügyelt eljárások minden konfigurációs beállításait megőrizve (15. táblázat), kizárólag így a tanulási halmazon változtatást (bővítést) eszközölve. A COS modellt futtattam az összes auditjelentésen, bemenetet képezve a tanulási és teszhalmaz számára is. A modell a lehetséges 115 kontrollból, 98 esetben tért vissza gyanúmomentummal, így a tanulási halmaz összesen 98 bináris változóval egészült ki, mely a COS modell ajánlásait tartalmazza. Mivel a felügyelt eljárások beállításai identikusak, ezért objektíven meg lehet győződni a hibrid megoldás javító/rontó hatásairól. A hibrid eljárás folyamatlépéseit a 25. ábra szemlélteti.



25. ábra: A hibrid modellezés folyamatábrája

Forrás: Saját szerkesztés

A hibrid modellt éleskörnyezetben futtatva a felügyelt módszerek performancia metrikái az alábbiak szerint alakultak (18., 19. és 20. táblázat). A hibrid megközelítést nélkülöző módszereket „egyszerű” modelleknek neveztem el.

18. táblázat: Egyszerű és hibrid ABM performancia metrikái

Performancia mutatók	Egyszerű ABM	Hibrid ABM	Változás mértéke (%)	Cél	Változás hatása (javulás/romlás)
<i>Igaz pozitívák száma (db)</i>	80	89	+11.25%	Növelés	Javulás
<i>Igaz negatívák száma (db)</i>	921	944	+2.50%	Növelés	Javulás
<i>Hamis pozitívák száma (db)</i>	42	19	-54.76%	Csökkentés	Javulás
<i>Hamis negatívák száma (db)</i>	79	70	-11.39%	Csökkentés	Javulás
<i>Pontosság</i>	0.89	0.92	+3.37%	Növelés	Javulás
<i>Precizitás</i>	0.66	0.82	+24.24%	Növelés	Javulás
<i>Fedés</i>	0.50	0.56	+12.00%	Növelés	Javulás
<i>F1-Pont</i>	0.57	0.67	+17.54%	Növelés	Javulás
<i>Variancia</i>	0.08	0.06	-0.25%	Csökkentés	Javulás
<i>AUROC</i>	0.85	0.90	+5.88%	Növelés	Javulás
<i>AUPRC</i>	0.58	0.70	+20.69%	Növelés	Javulás

Forrás: Saját szerkesztés

19. táblázat: Egyszerű és hibrid GBM performancia metrikái

Performancia mutatók	Egyszerű GBM	Hibrid GBM	Változás mértéke (%)	Cél	Változás hatása (javulás/romlás)
<i>Igaz pozitívák száma (db)</i>	46	62	+34.78%	Növelés	Javulás
<i>Igaz negatívák száma (db)</i>	954	958	+0.42%	Növelés	Javulás
<i>Hamis pozitívák száma (db)</i>	9	5	-44.44%	Csökkentés	Javulás
<i>Hamis negatívák száma (db)</i>	113	97	-14.16%	Csökkentés	Javulás
<i>Pontosság</i>	0.89	0.91	+2.25%	Növelés	Javulás
<i>Precizitás</i>	0.83	0.93	+12.05%	Növelés	Javulás
<i>Fedés</i>	0.29	0.39	+34.48%	Növelés	Javulás
<i>F1-Pont</i>	0.43	0.55	+27.91%	Növelés	Javulás
<i>Variancia</i>	0.04	0.03	-0.25%	Csökkentés	Javulás
<i>AUROC</i>	0.85	0.85	0%	Növelés	Semleges
<i>AUPRC</i>	0.61	0.71	+16.39%	Növelés	Javulás

Forrás: Saját szerkesztés

20. táblázat: Egyszerű és hibrid NN performancia metrikái

Performancia mutatók	Egyszerű NN	Hibrid NN	Változás mértéke (%)	Cél	Változás hatása (javulás/romlás)
<i>Igaz pozitívák száma (db)</i>	74	76	+2.70%	Növelés	Javulás
<i>Igaz negatívák száma (db)</i>	916	926	+1.09%	Növelés	Javulás
<i>Hamis pozitívák száma (db)</i>	47	37	-21.28%	Csökkentés	Javulás
<i>Hamis negatívák száma (db)</i>	85	83	-2.35%	Csökkentés	Javulás
<i>Pontosság</i>	0.88	0.89	+1.14%	Növelés	Javulás
<i>Precizitás</i>	0.61	0.68	+11.48%	Növelés	Javulás
<i>Fedés</i>	0.47	0.47	0%	Növelés	Semleges
<i>F1-Pont</i>	0.53	0.56	+5.66%	Növelés	Javulás
<i>Variancia</i>	0.11	0.11	0%	Csökkentés	Semleges
<i>AUROC</i>	0.82	0.85	+3.66%	Növelés	Javulás
<i>AUPRC</i>	0.56	0.62	+10.71%	Növelés	Javulás

Forrás: Saját szerkesztés

A hibrid modellezés szemmel láthatóan beváltotta a hozzá fűzött reményeket, szinte kivétel nélkül javulás mutatkozik az összes metrikában (minimálisan a GBM AUROC mutatója esetén is, azonban a kerekítés ezt nem engedi láttatni).

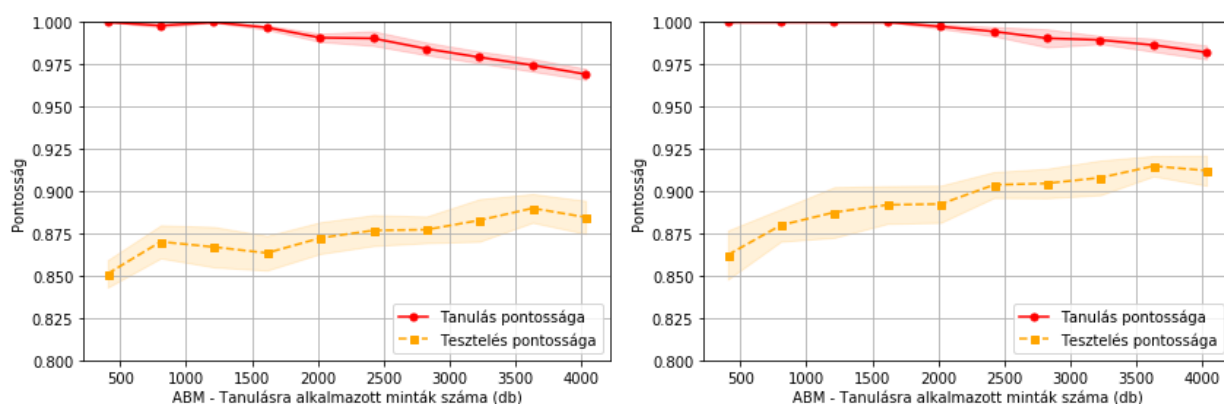
Az ABM algoritmus legnagyobb teljesítménynövekedése a hamis pozitív találatok jelentős csökkentésében érhető tetten, mely hatással van a Precizitás mutatóra. A hamis pozitívak efféle hangsúlyos redukálása az audit részéről erőforrás (humán és pénzügyi) megtakarítást enged, azaz kevesebb hamisan feltételezett hiányosságokkal rendelkező kontroll újraértékeléséről és ismételt teszteléséről szükséges meggyőződni. Kiemelendő még az igaz pozitív találatok, valamint AUROC és AUPRC mutatók növekedése, mely utóbbiak a legkedvezőbbek a három modell tekintetében.

A GBM eljárás igaz pozitív találatainak száma emelkedett drámaian, ezzel párhuzamosan a hamis pozitívak száma közel felére csökkent, mely az ABM-mel összehasonlítva is kicsivel több, mint negyede. A GBM erőssége így a minimális hamis pozitív találatok száma, illetve, a modell varianciája, mely jobb általánosító képességet feltételez, mely alapesetben kevésbé volt várható.

Az NN modell esetén is megfigyelhető a jóságmetrikák javulása, azonban kevésbé szembetűnő, mint az ABM és GBM modelleknél, mely valószínűsíthetően a modell eddig is feltárt (24. ábra) erős túlilleszkedésének köszönhető, mely a hibridizációt követően még inkább hangsúlyos.

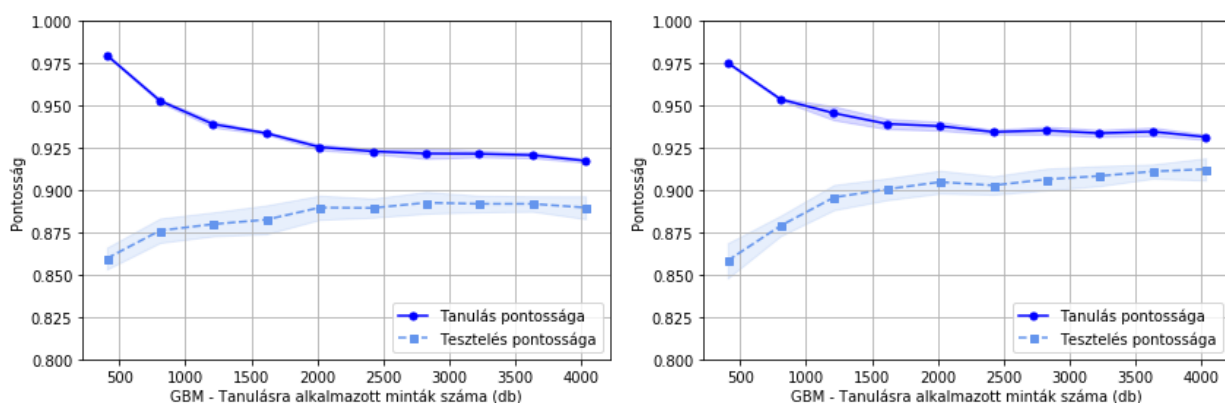
Általánosságban a hamis pozitív arány javult mind a három módszernél a leginkább, ezért a hibridizálásra alkalmazott COS modell, vélelmezhetően mérsékelte a modellek pozitív osztályban megtestesülő magabiztosságát és kevésbé „beszédesebb” eljárásokat hívott életre. A COS „ajánlásai” és a célváltozó értéke között így kimutatott összefüggés észlelhető, mely kihangsúlyozásra került az egymással leghasonlóbb objektumok között megfigyelt kontrollbeli eltérések révén, tehát a kontrollhiányosságok konstellációi súlyozottan kerültek feldolgozásra, melynek előnyei kézzel foghatóan tapasztalhatók.

A 4. és 5. melléklet tartalmazza modellenként az egyszerű és hibrid technikák ROC és PR-görbéit és a jóságmetrikák összehasonlító ábráit. A következő ábrák (26., 27., és 28. ábra) a hibrid alkalmazások tanulási görbéit szemléltetik, melyekben grafikusán is megállapítható a varianciák csökkenésének és a pontosság javulásának ténye (X tengely a tanulásra alkalmazott minták darabszáma, Y tengely a Pontosság mértéke 0-1 skálán).



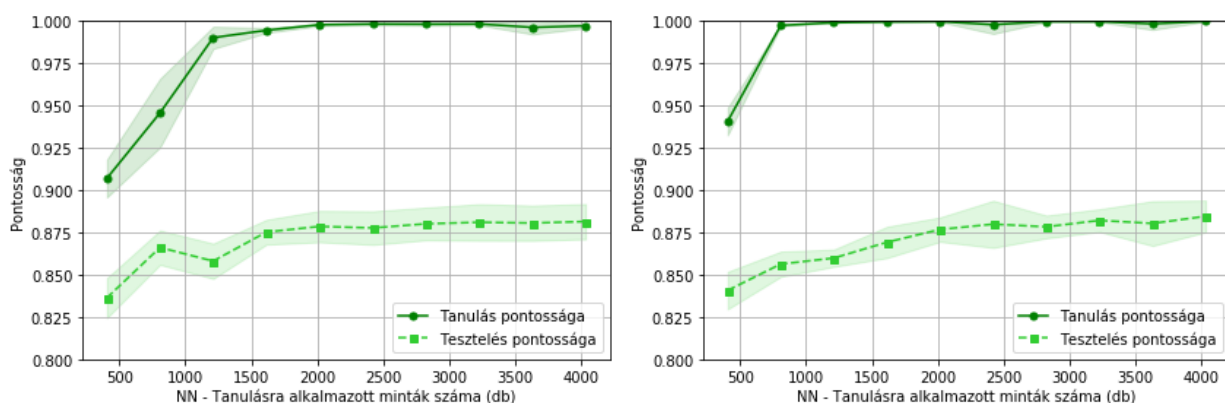
26. ábra: Az egyszerű és hibrid ABM algoritmus tanulási görbéi

Forrás: Saját szerkesztés



27. ábra: Az egyszerű és hibrid GBM algoritmus tanulási görbéi

Forrás: Saját szerkesztés



28. ábra: Az egyszerű és hibrid NN algoritmus tanulási görbéi

Forrás: Saját szerkesztés

4.2.7. Elemzések szűkített adathalmazokon

A terepmunka során gyűjtött adathalmaz a 4.1. alfejezetben ismertetett módon túlnyomórészt a pénzintézeti szektor szereplőire terjed ki (54%-a a teljes populációnak), valamint a könyvvizsgálathoz kapcsolódó informatikai vizsgálatokhoz köthetők (70%), így igény mutatkozik az adatvagyon kitüntetett részhalmazainak vizsgálatára az implementált gépi tanuló alkalmazásokat felhasználva. Azt szükséges vizsgálni, hogy:

- A teljes adathalmazon elért eredmények többségében csak a kitüntetett részhalmazok jelenlétének köszönhető-e;
- Szükséges-e külön döntéselőkészítő rendszert felépíteni az egyes részhalmazok számára, mert pontosabb előrejelzés várható az adatok homogén csoportosítása miatt.

Az első pont ellenőrzésére tekintsük meg a 21. és 22. táblázatot, melyek az egyszerű és hibrid felügyelt modellek performancia metrikáit részletezik audittípusonkénti bontásban (az egyszerűség és transzparencia kedvéért csak a Pontosság és F1-Pont mutatói kerültek feltüntetésre).

Mindkét táblázatból egyértelműen tetten érhető, hogy a könyvvizsgálathoz kapcsolódó informatikai vizsgálatok az átlagos, teljes adatvagyonra vonatkozó performancia metrikákat

meghaladják valamennyi modell esetén, tehát felhúzó erővel hatnak, mely a hibrid modellezés esetén még szembetűnőbb. Ez azt jelentheti, hogy az algoritmusok vagy elfogultak az audittípussal szemben, vagy elegendő adat áll rendelkezésre a bizonytalan döntések csökkentése érdekében, hogy annak általánosító képessége elfogadható legyen. Ennek tényéről akkor tudunk megbizonyosodni, ha elvégezzük a modellezést kizárólag a könyvvizsgálati kategóriában. Az is szemügyre vehető, hogy a jogszabályi megfeleléségi vizsgálatok minden esetben jóval az átlagos teljesítmény alatt maradtak, mely auditok kevésbe standard hatókörrel rendelkeznek, nem úgy, mint a könyvvizsgálathoz kapcsolódó informatikai vizsgálatok. A SOC auditok esetében a legkiemelkedőbb, hogy a kontrollok közötti összefüggéseket az egyszerű NN volt képes a legmagasabb mértékben leképezni, azonban a teljesítmény hibrid megközelítésben csökkent, azaz az ajánlórendszer nem volt képes hozzáadott értéket teremteni.

21. táblázat: Egyszerű felügyelt modellek performancia metrikái audittípusonkénti megoszlásban

Audittípusok	Egyszerű ABM		Egyszerű GBM		Egyszerű NN	
	Pontosság	F1-Pont	Pontosság	F1-Pont	Pontosság	F1-Pont
Könyvvizsgálathoz kapcsolódó informatikai vizsgálat	0.90	0.66	0.89	0.49	0.90	0.63
Egyéb	0.89	0.55	0.88	0.42	0.84	0.29
ISO 27001 audit/réselemzés	0.86	0.53	0.86	0.34	0.87	0.55
SOC/SOC2/SOC3 audit	0.92	0.50	0.92	0.36	0.93	0.61
Jogszabályi megfeleléségi vizsgálat	0.88	0.42	0.90	0.40	0.84	0.32
Teljes adatvagyon	0.89	0.57	0.89	0.43	0.88	0.53

Forrás: Saját szerkesztés

22. táblázat: Hibrid felügyelt modellek performancia metrikái audittípusonkénti megoszlásban

Audittípusok	Hibrid ABM		Hibrid GBM		Hibrid NN	
	Pontosság	F1-Pont	Pontosság	F1-Pont	Pontosság	F1-Pont
Könyvvizsgálathoz kapcsolódó informatikai vizsgálat	0.95	0.81	0.92	0.63	0.91	0.70
Egyéb	0.89	0.55	0.91	0.57	0.86	0.48
ISO 27001 audit/réselemzés	0.88	0.61	0.88	0.51	0.86	0.45
SOC/SOC2/SOC3 audit	0.93	0.48	0.92	0.36	0.92	0.44
Jogszabályi megfeleléségi vizsgálat	0.90	0.50	0.91	0.46	0.88	0.38
Teljes adatvagyon	0.92	0.67	0.91	0.55	0.89	0.56

Forrás: Saját szerkesztés

A 6. számú melléklet tartalmazza az egyszerű és hibrid felügyelt modellek performancia metrikáit iparági bontásban. Az egyes cellákban jelölt „n/a” azt jelenti, hogy adott szektorban az F1-Pont mutató értékét nem lehetett kiszámolni, mivel 0 került a nevezőbe, mely köszönhető annak, hogy nem volt vagy igaz pozitív, vagy hamis pozitív, vagy hamis negatív találat. Ez nem feltétlen jelenti azt, hogy az algoritmusok az adott iparágban egyöntetűen döntöttek minden minta esetében egy kategóriát célozva, hanem, hogy alapvetően a tesztalmazba nem került beválasztásra adott iparág

esetén kontrollhiányosság (ezt tapasztalható pl. az állami szféra mintáira). A táblázatból kiolvasható, hogy a pénzügyi szektor metrikái az egyszerű modellek esetén elmaradnak a teljes adatvagyonon mért mutatóktól, míg hibrid modellezésben átlag körüliek.

A modellezést (egyszerű és hibrid) elvégezve a két kiemelt kategóriára (könyvvizsgálathoz kapcsolódó informatikai vizsgálatok és pénzügyi szektor) a következő táblázatokban ismertetett eredményeket szolgáltatják az algoritmusok (23. és 24. táblázatok közvetítik a könyvvizsgálathoz kapcsolódó informatikai vizsgálatokra szűkített eredményeket, amíg ugyanezt a pénzügyi szektorra a 7. számú melléklet tartalmazza). A táblázatok bal hasábjában a kijelölt kategóriákat foglalja magában, a jobb hasábjában, a teljes adatvagyonon mért eredményeket ismertet, mely a kategóriára jellemző eredményekre korlátozódik. Mivel a két tanulómátrix különbözik, ezért a kategóriákra jellemző minták száma értelemszerűen eltér. A Variancia, AUROC és AUPRC mutató esetén a cellák értékei „n/a”, mivel a két mutató összesítve képes csak értéket szolgáltatni.

A táblázat alapján egyértelműen megállapítható, hogy a könyvvizsgálathoz kapcsolódó informatikai vizsgálatokon mért (szűkített adatvagyon) performancia metrikák, kizárólag a GBM algoritmus esetén haladják meg jelentősen a teljes adatvagyonon mért mutatókat, az ABM és NN módszerek esetén minimális eltérés tapasztalható.

A pénzügyi szektor adatait vizsgálva (7. számú melléklet) az egyszerű és hibrid modellek esetén is ideálisabb értékeket kaptunk a teljes adatvagyon felhasználásával, így egyértelmű, hogy a pénzügyi szektort nem érdemes elkülöníteni, a többi szektor leíró adata segíti a döntéshozatást. A könyvvizsgálathoz kapcsolódó informatikai vizsgálat esetén az eddigiek alapján úgy tűnik, hogy érdemes lehet külön döntéshozatást rendszer fenntartása a GBM algoritmussal, amennyiben a vezetői döntés annak implementálása (nem számolva az üzemeltetéshez kapcsolódó költségeket).

23. táblázat: Felügyelt egyszerű módszerek performancia metrikái szűkített és a teljes adatvagyonon

Performancia mutatók	Szűkített adatvagyon (könyvvizsgálathoz kapcsolódó informatikai vizsgálatok)			Teljes adatvagyon (csak könyvvizsgálathoz kapcsolódó informatikai vizsgálatok eredményei)		
	Egyszerű ABM	Egyszerű GBM	Egyszerű NN	Egyszerű ABM	Egyszerű GBM	Egyszerű NN
Igaz pozitívák száma (db)	44	43	40	43	24	37
Igaz negatívák száma (db)	337	343	346	355	370	362
Hamis pozitívák száma (db)	22	16	13	22	7	15
Hamis negatívák száma (db)	29	30	33	23	42	29
Pontosság	0.88	0.89	0.89	0.90	0.89	0.90
Precizitás	0.67	0.73	0.75	0.66	0.77	0.71
Fedés	0.60	0.59	0.65	0.65	0.36	0.56
F1-Pont	0.63	0.65	0.63	0.66	0.49	0.63
Variancia	0.09	0.03	0.08	n/a	n/a	n/a
AUROC	0.88	0.90	0.91	n/a	n/a	n/a
AUPRC	0.71	0.75	0.74	n/a	n/a	n/a

Forrás: Saját szerkesztés

24. táblázat: Felügyelt hibrid módszerek performancia metrikái szűkített és a teljes adatvagyonon

Performancia mutatók	Szűkített adatvagyon (könyvvizsgálathoz kapcsolódó informatikai vizsgálatok)			Teljes adatvagyon (csak könyvvizsgálathoz kapcsolódó informatikai vizsgálatok eredményei)		
	Hibrid ABM	Hibrid GBM	Hibrid NN	Hibrid ABM	Hibrid GBM	Hibrid NN
<i>Igaz pozitívák száma (db)</i>	53	47	43	49	32	45
<i>Igaz negatívák száma (db)</i>	353	356	348	371	374	359
<i>Hamis pozitívák száma (db)</i>	6	3	11	6	3	18
<i>Hamis negatívák száma (db)</i>	20	26	30	17	34	21
<i>Pontosság</i>	0.94	0.93	0.91	0.95	0.92	0.91
<i>Precizitás</i>	0.90	0.94	0.80	0.89	0.91	0.71
<i>Fedés</i>	0.73	0.64	0.59	0.74	0.48	0.68
<i>F1-Pont</i>	0.80	0.76	0.68	0.81	0.63	0.70
<i>Variancia</i>	0.06	0.02	0.07	n/a	n/a	n/a
<i>AUROC</i>	0.91	0.93	0.90	n/a	n/a	n/a
<i>AUPRC</i>	0.84	0.88	0.78	n/a	n/a	n/a

Forrás: Saját szerkesztés

A modellekhez tartozó ROC és PR-görbéket, valamint a tanulási görbéket a 8. és 9. számú mellékletek tartalmazzák. Kiemelendő, hogy mind a két kategóriában az ABM algoritmus (piros és narancssárga színnel jelölve) erős túlilleszkedést produkált, mivel a tanuló halmazon mért pontossági adatok közel 100%-osak már az első 100 db minta feldolgozása után. Továbbá, általánosságban elmondható, hogy a tesztelési halmazon keresztvalidációs eljárással mért eredmények varianciái magasak, tehát instabil modellekről lévén szó, összességében az alfejezetben leírtakkal összhangban, véleményem szerint, nem javasolt külön döntéstámogató rendszerek létrehozása a két részleteiben is megvizsgált kategóriák egyikében sem.

4.2.8. Objektív modelljóság-becslés a generált modelleken anti-diszkriminatív eljárással

A generált modellek kiértékelése és a legjobb modell meghatározása ez idáig szubjektív kereteken belül történt, melyet a döntéshozói preferenciához kötöttem. Belátható, hogy különböző erőforrással és kockázati étvággyal rendelkező szervezetek eltérően vélekedhetnek a megoldások alkalmazhatóságáról, például, a GBM modell bizonyult a legstabilabb modellnek (alacsony variancia), az ABM esetén volt a megmagasabb az F1-Pont értéke, az NN performancia metrikái pedig a legtöbb esetben a másik két modell mutatói között helyezkedtek el, így jogosan felmerül a kérdés, hogy egy robot döntéshozó, melyik modellt választaná.

Az anti-diszkriminatív modellezés képes a kérdésre választ adni. Felügyelet nélküli modellezés keretében a hasonlóságelemzés matematikai apparátusát felhasználva egy fiktív célváltozó becslése szükséges, mely képes a normától történő elmozdulás mértékének meghatározása optimalizáltan, kihasználva a függvény-szimmetria által adott lehetőségeket. A hasonlóságelemzés elvégzéséhez szükséges az egyes modellek közötti teljesítménymutatók alapján rangsort felállítani, ahol a bemenetként felhasznált relativizált metrikák (Pontosság, Precizitás, Fedés, F1-Pont, Variancia, AUROC, AUPRC) bármely modell esetén biztosítják az összehasonlíthatóságot. A rangsorolt adatokra lépcsős függvényt illesztve a 25. táblázatban szemléltetett eredményeket rögzítettem, ahol az Y_0 oszlop a fiktív célváltozó becsült értékét

ismerteti, ahol a tényadat (norma) 1000 volt. Összesen 18 modell került összehasonlításra, ahol a „TA” a teljes adatvagyonon, a „KKIV” a könyvvizsgálathoz kapcsolódó informatikai vizsgálatokra, illetve a „PSZ” a pénzügyi szervezetre szűkített adatbázison futtatott modelleket jelöli.

25. táblázat: Modelljóság-becslés anti-diszkriminatív eljárással

Modellek	Pontosság	Precizitás	Fedés	F1-Pont	Variancia	AUROC	AUPRC	Y₀
<i>Egyszerű ABM - TA</i>	0.89	0.66	0.50	0.57	0.08	0.85	0.58	997.50
<i>Egyszerű GBM - TA</i>	0.89	0.83	0.29	0.43	0.04	0.85	0.61	972.90
<i>Egyszerű NN - TA</i>	0.88	0.61	0.47	0.53	0.11	0.82	0.56	961.90
<i>Hibrid ABM - TA</i>	0.92	0.82	0.56	0.67	0.06	0.90	0.70	1027.60
<i>Hibrid GBM - TA</i>	0.91	0.93	0.39	0.55	0.03	0.85	0.71	1021.60
<i>Hibrid NN - TA</i>	0.89	0.68	0.47	0.56	0.11	0.85	0.62	980.90
<i>Egyszerű ABM - KKIV</i>	0.88	0.67	0.60	0.63	0.09	0.88	0.71	1003.50
<i>Egyszerű GBM - KKIV</i>	0.89	0.73	0.59	0.65	0.03	0.90	0.75	1027.60
<i>Egyszerű NN - KKIV</i>	0.89	0.75	0.65	0.63	0.08	0.91	0.74	1026.60
<i>Hibrid ABM - KKIV</i>	0.94	0.90	0.73	0.80	0.06	0.91	0.84	1052.70
<i>Hibrid GBM - KKIV</i>	0.93	0.94	0.64	0.76	0.02	0.93	0.88	1055.70
<i>Hibrid NN - KKIV</i>	0.91	0.80	0.59	0.68	0.07	0.90	0.78	1031.60
<i>Egyszerű ABM - PSZ</i>	0.88	0.57	0.45	0.50	0.10	0.78	0.49	956.90
<i>Egyszerű GBM - PSZ</i>	0.90	0.78	0.32	0.46	0.05	0.79	0.52	984.00
<i>Egyszerű NN - PSZ</i>	0.87	0.51	0.35	0.42	0.09	0.75	0.38	950.40
<i>Hibrid ABM - PSZ</i>	0.90	0.73	0.45	0.55	0.08	0.80	0.55	990.50
<i>Hibrid GBM - PSZ</i>	0.91	0.89	0.38	0.54	0.04	0.82	0.63	1007.50
<i>Hibrid NN - PSZ</i>	0.87	0.54	0.32	0.40	0.09	0.76	0.40	950.90

Forrás: Saját szerkesztés

A táblázat alapján megállapítható, hogy összesítve a Hibrid GBM – KKIV adathalmazon mért mutatói bizonyulnak a legideálisabbnak (1055.70) a felhasznált metrikák tükrében. A teljes adatvagyonon tanult módszerek közül egyszerű és hibrid kategóriában is kiemelkedik az ABM, olyannyira, hogy jobb értéket produkált az egyszerű ABM (997.50) a hibrid NN-nél is (980.90). Az említett kategóriák második helyezettjei a GBM, a harmadik az NN algoritmusok. Az ABM győzelme azt engedi vélelmezni, hogy relatív sok volt a kiugró érték az adathalmazban, mivel az ABM az adatvagyon ismételt súlyozását végzi az egyes iterációkban.

A KKIV adathalmazon a GBM algoritmusok teljesítettek a legjobban (1027.60 és 1055.70), melyet az egyszerű modellek esetén az NN (1026.60), hibrid modellek esetén az ABM (1052.70) követ, tehát a hibrid megközelítés nem egységes eloszlásban javította a modellek predikációs képességét.

A PSZ adathalmazon, hasonlóan a GBM algoritmusok mutatták a legideálisabb értékeket (984.00 és 1007.50), melyet rendre az ABM (956.90 és 990.50) és NN (950.40 és 950.90) módszerek követnek egyszerű és hibrid modellezés esetén is.

Kijelenthető, hogy objektív eljárással vizsgálva, a teljes adatvagyonon a döntéstámogató rendszer az ABM eljárást, a könyvvizsgálathoz kapcsolódó informatikai vizsgálatokra és pénzügyi

szektorra korlátozott halmazokon a GBM modelleket választaná, amennyiben a modellszelekció is automatizált döntéselőkészítéshez lenne kötve.

Megvizsgálva a 25. táblázatban közölt modellek leíró tulajdonságai alapján származtatott csoportokat (kategóriákat), a 26. táblázat ismerteti a modell-idealitások (Y_0) átlagait, melyre a csoportátlagok eltérésének vizsgálatára volt szükség. Annak ellenőrzésére, hogy a kategóriaátlagok szignifikánsan eltérnek-e, varianciaelemzést alkalmaztam melyet a 10. számú melléklet részletez. A kutatói kérdés: a kategóriagyőztesek alátámasztják-e a Hibrid GBM – KKIV első helyezését?

26. táblázat: Modell-jóság átlagok kategóriánként

Csoportosítás	Csoportok	Y_0 átlagok
<i>Modellkomplexitás</i>	Egyszerű modellek	986.81
	Hibrid modellek	1013.22
<i>Alkalmazott módszer</i>	ABM	1004.78
	GBM	1011.55
	NN	983,72
<i>Felhasznált adatvagyon</i>	TA	993.73
	KKIV	1032.95
	PSZ	973.37

Forrás: Saját szerkesztés

A táblázatból kiolvasható, hogy a hibrid modellek, átlagosan jobbak voltak (1013.22) az egyszerű modelleknél (986.81), mely a korábban bemutatott eredmények alapján elvárt volt. Az eltérés nem tekinthető szignifikánsnak (0.10) a 10. számú melléklet alapján, mivel az F-próba szignifikanciaszintje meghaladja a társadalomtudományokban általánosan elvárt 0.05-os szintet².

A GBM modell (1011.55) átlagosan jobban teljesített az ABM (1004.78) és NN (983.72) modelleknél, a különbség az előző kategóriához képest hasonlóan nem tekinthető szignifikánsnak (0.35) a 10. számú melléklet alapján.

A KKIV (1032.95) adathalmazon mért teljesítmények szignifikánsan (0.00) jobbak (10. számú melléklet), mint a TA (993.73), valamint PSZ (973.37) adatvagyonokon mért mutatók, melyet a Scheffé-próba³ is alátámaszt: a KKIV adathalmazon mért eredmények szignifikánsan különböznek, a TA és PSZ között nincs szignifikáns különbség.

Kijelenthető, hogy csoportosítva a modelleket, a Hibrid GBM – KKIV modell mindegyik kategóriában a legideálisabb leíró tulajdonsággal rendelkezik, így valóban a legjobb modellnek tekinthető egy új konzisztencia-réteg megerősítése alapján is, mely automatizáltan került megállapításra.

² A 10. számú melléklet, többek között, részletezi a Levene-teszt szignifikanciáját, mely értékekből megállapítható, hogy a szóráshomogenitás feltétele teljesül. A Levene-teszt nullhipotézise: a szórások nem egyenlők, elvetésre kerül.

³ A Scheffé-próba az egyes kategóriák közötti páros összehasonlítások szignifikanciáját ismerteti, melyet csak 2-nél több kategóriára lehet elvégezni, melyek között volt szignifikáns (F-próba) különbség.

4.2.9. Az alfejezet összefoglalása

Az alfejezetben bemutatásra került a terepmunkán gyűjtött adatvagyon felügyelt és felügyelet nélküli, valamint hibrid gépi tanuló módszerekkel történő feldolgozása és részletes kiértékelése. A közöltekkel kapcsolatban, összefoglalóan, az alábbiak állapíthatók meg:

- Objektíven bizonyítást nyert, hogy a kontrollhiányosságok megléte és azok konstellációi összefüggéseket tartalmaznak, osztályozási problémaként történő megközelítésben az alkalmazott algoritmusok képesek voltak a véletlen találgatásnál ideálisabb eredményt szolgáltatni, azaz előrejelezni a gyanúk meglétét;
- A predikció sikeresnek tekinthető, mind az alkalmazott felügyelt, mind a felügyelet nélküli algoritmusok alkalmazásával;
- A felügyelt és felügyelet nélküli módszerek összehangolt hibrid felhasználásával az idealitás mérésére szolgáló performancia metrikák egyértelműen jobban teljesítettek, mint egyszerű megközelítésben. Ez köszönhető annak, hogy a felügyelet nélküli ajánlórendszer hasonlóság-alapon képes volt helyesen meghatározni azon kontrollokat, ahol „illendő” volt a hiányosságok visszaigazolása, ezzel addicionális többlet információt teremtve, súlyozva a megállapítások tényét, felerősítve a kontrollok közötti interakciók jelenlétét;
- A hibrid modellezés keretében az ABM és GBM algoritmusok varianciái ideálisabb értéket vettek fel, az NN esetén nem történt sem javulás, sem romlás;
- Külön vizsgálatra került a teljes adatvagyon két jelentős részhalmaza, mely alapján megállapítható, hogy nem érdemes szeparált döntéstámogatórendszer fenntartása a specifikusabb auditok számára.

A leírtak alapján kijelenthető, hogy az alfejezet elején ismertetett hipotézisek igazolhatók.

4.3. Genetikai potenciál keresése a gépi tanulás adatvagyonának redukált felhasználásával

A gépi tanuló módszerek legnagyobb kihívásai közé sorolható a rendelkezésre álló tanulási halmazból történő maximális tudás kinyerése, mely alapján egy éles üzemi környezetben az alkalmazott rendszer képes még nem látott adatokon elvégezni a döntéstámogatást, tehát elfogadható/maximális általánosító képességekkel rendelkezik. A tanuló minta gyakran zajos, nem releváns (ill. véletlen-szerű) információt tartalmazhat, ezért a szakirodalmi áttekintésben is ismertetett módszerek alkalmazhatók, azonban az alfejezet egy újfajta kísérletet és bizonyítást szándékozik bemutatni, melynek központi eleme a tanulási halmaz redukálása, azaz azon rekordok elhagyása, melyek a legkevesebb értékes vagy leginkább kockázatos információt hordozzák a döntéselőkészítés jóságmetrikáinak növeléséhez, illetve visszatartó erővel rendelkeznek, megtévesztik a felhasznált algoritmusokat.

Az inkonzisztens/kockázatos rekordok eredhetnek emberi és szakmai tévedésből, az adatrögzítés nem megfelelőségéből, vagy az adattanszformációs eljárások hibás alkalmazásából - s leginkább a mérés-definícióból magából – hiszen az emberi önkényes állapotmeghatározások nem mérések – a fizikai valóságról alkotott ismételtetőségi elvet követő, a mérési hibát értelmezni engedő mérés-fogalommal szemben. A tanulási halmaz optimalizált felhasználásával a probléma megoldására irányuló gépi tanuló algoritmusok genetikai potenciálja kiaknázható.

Jelen alfejezet az előző alfejezetben fejlesztett algoritmusok a tanulóhalmaz optimálisához közelítő felhasználásának (a genetikai algoritmusokhoz hasonlóan speciálisan) intelligens keresési eljárással történő javítását célozza meg, az alábbi hipotézis igazolását ismerteti:

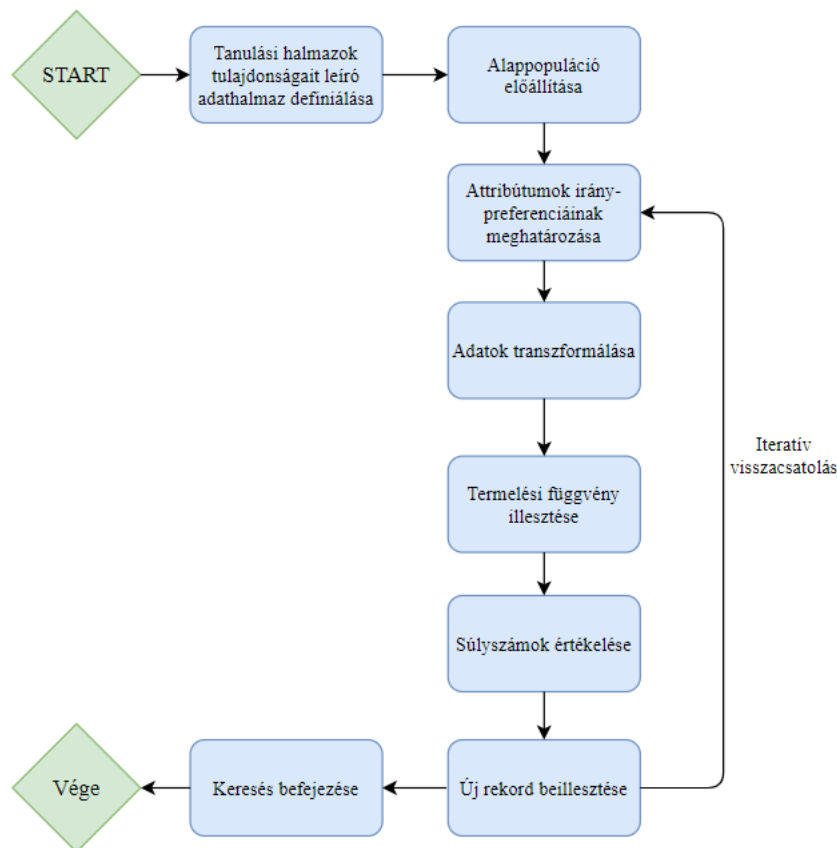
H2: A döntéstámogató rendszer genetikai potenciálja letapogatható hasonlóságelemzéssel ellátott kereső eljárással a tanításra alkalmazott adathalmaz irányított feldolgozásán keresztül, úgy, hogy a genetikai potenciálhoz vezető kereső eljárás a genetikus algoritmusok esetén alkalmazott véletlen mutáció és a populáció egyedeinek keresztezése nélkül is képes ideálisabb eredményt szolgáltatni.

4.3.1. A javasolt innovatív keresési eljárás

A tanuló algoritmus tanulási halmazának optimális felhasználásával meghatározott genetikai potenciálja keresési algoritmussal közelítendő, melyben a keresés akkor járt sikerrel, ha a kiindulóponttól különböző, objektíven felismerhető, ideálisabb eredményt érünk el, egy vagy több lépésben úgy, hogy a genetikai potenciál, mint szélsőérték azonosítható be iteratív módon. A keresés irányának meghatározására a hasonlóságelemzés lépcsős függvényei megfelelő matematikai apparátust kínálnak, ahol termelési függvények illesztésével és felhasználásával szükséges/lehetséges a dedikált célváltozó elmozdítása a kívánt irányba, a rendelkezésre álló attribútumok érték-irány növelésével/csökkentésével, tehát a független változók tudatos változtatásával iteratív módon. Az attribútumok a tanulásra felhasznált adathalmaz tulajdonságait kell, hogy szükségszerűen leírják, míg az objektumok az alternatív megoldások egy tetszőleges/tudatos halmazát jelentik.

A termelési függvények lépcsős függvényeket alkalmaznak a súlyozásra, így az iteráció akkor fejeződik be, ha a lépcsős függvények legelőkelőbb súlyszámainak összege megegyező vagy kisebb értéket mutatnak az aktuális célváltozó értékénél, vagy pl. a polinomhatás következtében

belátható, hogy a súlyszámok maximuma nem elérhető. Lépcsős függvények generálásához szükséges egy (véletlen vagy tudatos) alappopuláció meghatározása, mely az alternatívák között rangsort felállítva, az elemzés lépéseként képes az elmozdulás irányát optimálisan definiálni, a lehető legjobb értékek maximalizálásával. A keresés kiindulópontja az eddigi legjobb eredmény (célváltozó érték). A leírt algoritmus értelmezhető egy speciális „genetikus” (innovatív) kereső eljárásnak, azonban az egyes iterációk között nem került alkalmazásra mutáció, azaz véletlen adatmanipuláció, valamint a populáció tagjainak keresztezése, és nincs szükség minden generációban sok-sok véletlen objektum kialakítására sem a kezdeti populáción túl. Az eljárás menetét vizuálisan szemlélteti a következő ábra (29. ábra).



29. ábra: A javasolt, a genetikai potenciált kereső eljárás folyamatábrája

Forrás: Saját szerkesztés

A genetikai potenciál keresésére javasolt algoritmus az alábbiakban leírt lépésekben képes az ideálisabb célváltozó heurisztikájára:

- 1. Tanulási halmazok tulajdonságait leíró adathalmaz definiálása:** Határozzuk meg azon attribútumokat, melyek vélhetően alkalmasak a tanulóhalmaz tulajdonságainak objektív leírására, valamint azon feltételeket, melyek mentén megtörténik a tanulóhalmaz redukálása. Az eljárás eredményességéhez szükséges olyan attribútumok előállítása, melyek képesek a célváltozó jóságához hozzájárulni. Az attribútum részét képezi, tehát, a célváltozó értékek, melyek egy adott tanuló halmazon történő algoritmus futtatásához köthető (annak egy tanulási rekordra jutó performanciáját írja le).

2. **Alappopuláció előállítás:** A definiált feltételek mentén töltsük fel az adattáblát (független változók és függő változó), ahol minden egyes rekord egy gépi tanuló algoritmus futtatásának felel meg. Ez a feltöltés lehet véletlenszerű és lehet a leíró mutatók által kijelölt kombinatorikai tér alapján szisztematikus is.

A 3. – 7. pontok iterációja, amíg az elfogadható célváltozó értéket meg nem találjuk:

3. **Attribútumok irány-preferenciáinak meghatározása:** Vizsgáljuk meg a független és függő változók kapcsolatát Pearson-féle korrelációval, és a korrelációs együttható előjelének függvényében határozzuk meg minden egyes attribútumra annak irány-preferenciáját. Természetesen, az irány-preferenciák levezetésére számos más matematikai lehetőség is adódik, mint pl. a hasonlóságelemzés is alkalmazható (lásd 4.4.5. alfejezet), azonban az egyszerű és gyors számítás kedvéért az eljárásban a Pearson-féle korrelációt használtam. Ideális esetben az attribútum definíciójából szervesen következik annak iránya is.
4. **Adatok transzformálása:** Rangsoroljuk az irány-preferenciák mentén az adathalmaz független változóit.
5. **Termelési függvény illesztése:** A rendelkezésre álló populációra illesszünk termelési függvényt.
6. **Súlyszámok értékelése:** A termelési függvény által előállított súlyszámokat kiértékelve, határozzuk meg, hogy az eddig (maximális) legideálisabbnak ítélt célváltozóval rendelkező rekord mely attribútumát szükséges változtatni a még ideálisabb eredmény reményében (ahol a legnagyobb az aktuális és következő súlyszám különbsége), tehát milyen típusú és tulajdonságokkal rendelkező auditjelentést szükséges a tanulóhalmazból kivonni. Amennyiben az algoritmus azt kívánja, hogy már előállított tanulási halmazt hozzunk ismételtlen létre, a változtatni kívánt attribútum szerinti súlyszámot számszerűen megugorva/túllépve, egy soron következő auditjelentést is vonjunk el a tanulóhalmazból (duplikált feldolgozást kiküszöbölve). Ha a legideálisabb súlyszámok összege kevesebb, vagy egyenlő, mint a jelenlegi célváltozó értéke, állítsuk meg a keresést.
7. **Új rekord beillesztése:** Futtassuk az algoritmust, az új rekordot szűrjük be az adattáblába. Amennyiben jobb eredményt érünk el, az új rekord lesz a keresés új kiindulópontja, azaz benchmark értéke. Abban az esetben, ha az eredmény kedvezőtlenebb, mint a kiindulópont, akkor az változatlan marad és szükséges lehet az irány-preferenciák felülbírlata (a 3. pontban leírtakkal összhangban).
8. **Keresés befejezése:** A kilépési feltételek teljesülésével a leideálisabb rekordot hirdessük ki a keresés győztesének.

Az algoritmus pszeudokódja megtalálható a 11. számú mellékletben.

Az algoritmus az alábbi korlátokkal alkalmazható a genetikai potenciál keresésére:

- Követelmény a racionális attribútumok létrehozása, melyek nélkül a keresési eljárás hibás összefüggéseket „tárhat fel”;
- Az eljárásban nincs mutáció, ezért az alappopuláció meghatározása hatékonysági szempontból kritikus, mivel annak minőségétől függ a keresés gyorsasága. Előfordulhat, hogy az alappopuláció már egy olyan rekordot is tartalmaz, mely a termelési függvény szerint a genetikai potenciál csúcsa, amit azonnal ki kell és ki lehet tudni mutatni;
- Lehetséges, hogy egy adott attribútumnak a valóságban optimuma van (azaz monoton nem irányítható), ezért a súlyszámok összegeként becsült genetikai potenciál csak elméleti;
- Az attribútumok összefüggését a célváltozóval folyamatosan ellenőrizni és konfirmálni szükséges, mert a meghatározott populáció csak egy mintája a teljes kombinatorikai térben elhelyezhető pontoknak és közel sem tekinthető reprezentatívnak.

4.3.2. Genetikai potenciál keresése a rögzített adathalmazon

Tanulási halmazok tulajdonságait leíró adathalmaz definiálása

Az adattábla előállításához szükséges attribútumokat kvázi véletlenszerűen határoztam meg, törekedve azon követelmény teljesítésére, hogy az adattáblák legyenek informatívak és képesek arra, hogy az egyes tanuló halmazok főbb tulajdonságait minimális redukálással is rekordonként változó értékekkel írják le. Továbbá, szem előtt tartottam azon, az eredeti adathalmazt leíró attribútumokat is, melyek a legnagyobb mértékben reprezentálják az auditjelentések sokaságát, bízva, hogy képesek hozzájárulni a célváltozó értékéhez. Az alábbi attribútumokat határoztam meg a kísérletben:

- f_1 - Átlagos megállapítások száma
- f_2 - Megállapítások szórása
- f_3 - Átlagos kontrollok száma, ahol nem volt megállapítás
- f_4 - Azon kontrollok szórása auditjelentéseként, ahol nem volt megállapítás
- f_5 - Pénzügyi szektorban vizsgált auditjelentések száma
- f_6 - Energetikai szektorban vizsgált auditjelentések száma
- f_7 - Könyvvizsgálathoz kapcsolódó informatikai vizsgálatok száma
- f_8 - Jogszabályi megfeleléségi vizsgálatok száma
- f_9 - Hozzáférs szabályozás területén azonosított átlagos megállapítások száma
- f_{10} - Működési biztonság területén azonosított átlagos megállapítások száma
- Y – F1-Pont

A maximalizálandó célváltozóknak az F1-Pontot jelöltem ki, mely a Fedés és Precizitás harmóniáját szemlélteti, azonban fel lehetett volna használni, például a Pontosságot, ARUOC értéket és akár az Igaz pozitív találatok számát, illetve komplexebb értelmezésben párhuzamosan ezek mindegyikét is. Az F1-Pont 0 és 1 közötti értékeket vesz fel, azonban ezt a számot megszoroztam 1000-el, mert a hasonlóságelemzésre alkalmazott online környezet kizárólag pozitív számokkal tér vissza, és a futtatás első fázisában nem ismert, milyen mértékben becsüli alá/fölé a termelési függvény a legkisebb /legnagyobb értékeket a függő változók esetén.

Alappopuláció előállítás

A tanulási halmaz redukálására 10 feltételt határoztam meg, melyből rendezve az auditjelentéseket a felső 5 és 10 auditjelentést lépésenként elvontam, tehát az alappopulációban összesen 21 rekord szerepel, melyből a legelső a teljes tanulási halmazon mért adatokat prezentálja. A feltételek az alábbiak:

- Legkevesebb megállapítással rendelkező auditjelentések
- Legtöbb megállapítással rendelkező auditjelentések
- Legkisebb hatókörrel rendelkező auditjelentések
- Legnagyobb hatókörrel rendelkező auditjelentések
- Legkisebb megállapítás aránnyal rendelkező auditjelentések (megállapítás / nincs megállapítás)
- Legnagyobb megállapítás aránnyal rendelkező auditjelentések (megállapítás / nincs megállapítás)
- Legkisebb megállapítás szórással rendelkező auditjelentések
- Legnagyobb megállapítás szórással rendelkező auditjelentések
- Pénzügyi szektor legkevesebb megállapításai
- Pénzügyi szektor legtöbb megállapításai

A feltételek mentén redukált és az attribútumok által leírt tanulólthalmazok adatait tartalmazó adattáblát a 12. számú melléklet ismerteti.

A gyakorlati szemléltetés példájának a hibrid GBM algoritmust választottam, azonban az alfejezet végén a hibrid ABM és hibrid NN eredmények is közlésre kerülnek. Az algoritmusok konfigurációit nem változtattam, kizárólag a tanulási halmaz összetételét (ceteris paribus), ezért a korábban bemutatott eredményekhez képest objektíven összehasonlíthatók a jóságmutatók.

A keresés iterációi

Az iteráció első lépése az irány-preferenciák meghatározása, mely a Pearson-féle korreláció együtthatójának előjeléből levezethető az egyes attribútumok és a célváltozó között. A korrelációs értékeket az alábbi táblázat szolgáltatja (27. táblázat). Az f_5 , f_6 , f_7 és f_8 az auditjelentések számával áll összefüggésben, minél kisebb az auditjelentések száma, annál kevesebb lesz az attribútumok értéke értelemszerűen. Mivel a cél az optimális tanulólthalmaz megtalálása a tanulási pontok redukálása által és nincs lehetőség további adatok gyűjtésére, így racionális az a döntés, hogy bár pozitív korreláció vélhető a hivatkozott attribútumok és célváltozó között, mesterségesen szükséges azokat megváltoztatni, melyet a 27. táblázat utolsó oszlopa szemléltet.

Az automatizálhatóság kapcsán tehát először a definitív irányokat kell meghatározni, majd ezek hiányában a korrelációk előjelét preferálni.

27. táblázat: Az alappopuláció irány-preferenciáinak meghatározása

Attribútumok	Pearson-korreláció (f_x, Y)	Irány-preferencia	Módosított irány-preferencia
f_1	0.46	Minél nagyobb, annál jobb	n/a
f_2	0.36	Minél nagyobb, annál jobb	n/a
f_3	-0.40	Minél kisebb, annál jobb	n/a
f_4	-0.34	Minél kisebb, annál jobb	n/a
f_5	0.37	Minél nagyobb, annál jobb	Minél kisebb, annál jobb
f_6	0.17	Minél nagyobb, annál jobb	Minél kisebb, annál jobb
f_7	0.04	Minél nagyobb, annál jobb	Minél kisebb, annál jobb
f_8	0.55	Minél nagyobb, annál jobb	Minél kisebb, annál jobb
f_9	0.79	Minél nagyobb, annál jobb	n/a
f_{10}	0.44	Minél nagyobb, annál jobb	n/a

Forrás: Saját szerkesztés

Az irány-preferenciák meghatározásával az adatok rangsorolása a feladat. A rangsorolt adathalmazt a 13. számú melléklet részletezi, melyből egyértelműen látszik, hogy az alappopuláció o9 azonosítóval ellátott objektuma lesz a kiindulópont, mely az előzetesen mért F1-Pont értékét meghaladja 0.58-dal, tehát már az alappopuláció inicializálásakor ideálisabb eredményt sikerült elérni a mindenkor konstans teszhalmazon a tanulóhalmaz csökkentésével. A 12. és 13. számú melléklet táblázatait áttekintve felmerülhet a kérdés, miért lehetséges, hogy két azonos értékkel rendelkező objektum rangsorszáma eltérő. A magyarázat a megjelenített értékek kerekítésében keresendő, a táblázatban egy tizedesjegyig kerekítve szerepelnek a számok, valójában a kettő között különbség van.

A rangsorolt adathalmazra így termelési függvény illeszthető a hasonlóságelemzés speciális algoritmusát alkalmazva, mely minden egyes attribútumhoz kalkulál egy lépcsősfüggvényt. A lépcsősfüggvény értékei attribútumonként az alábbi (28. táblázat), ahol kijelölésre került halványzölddel az o9 jelenlegi súlyszámai, valamint sötétzölddel a következő lépcső, melyet szükségszerűen az attribútumok megváltoztatásával el kell érni (ahol csak halványzöld szerepel, ott az o9 jelenleg maximális és nem tud jobb súlyszámot elérni a jelen populációban).

A táblázat alapján kiszámolható, hogy a következő legnagyobb súlyszámot az f_{10} irány-preferencia szerinti növelésével lehet elérni, ahol annak értéke 1.9 (12. számú melléklet o16 azonosítóval ellátott objektuma), tehát, olyan auditjelentést szükséges elvenni, ahol kevés megállapítás szerepel az A12. A működés biztonsága kontrollterületen. Az is látható a táblázatból, hogy a legnagyobb javulást, elméletileg, elérni az f_8 attribútum első súlyszámának megugrásával lehetséges.

A súlyok együttes kerekített értéke 1030, amely túlmutat az elméletileg lehetséges genetikai potenciálon, mivel az F1-Pont soha nem lehet 1-nél nagyobb szám, azonban azt szemlélteti az érték, hogy van még potenciál a tanulásban, ezért érdemes folytatni az iterációt. Kiemelendő, hogy a modellben semmi nem tanítja a rendszert az elméleti maximum tényének levezetésére, de természetesen, ha ilyen létezik, a megszorítást erre a maximumra bármikor aktiválni lehet. Azonban az ideális állapot értelmében a maximum alatti genetikai potenciál esetét kell, hogy közelítse több lépésben a keresés – így lényegtelen, mekkora az amplitúdó az egyensúlyi végérték körül. Továbbá, a táblából az is látszik, hogy a jelen alappopulációban a teljesítmény növelésében az f_7 egyelőre nem játszik szerepet, az attribútum és a célváltozó között az első iterációban rendelkezésre álló adatok alapján nincs összefüggés.

28. táblázat: Az alappopuláció súlyszámai az első iterációban

Lépcsők	f ₁	f ₂	f ₃	f ₄	f ₅	f ₆	f ₇	f ₈	f ₉	f ₁₀
S1	103.5	3.5	122.5	182	42	209	0	179.5	69	118.5
S2	103.5	0	122.5	182	42	202	0	169	69	118.5
S3	103.5	0	122.5	159.5	37.5	0	0	28	69	101
S4	103.5	0	111.5	89.5	37.5	0	0	28	69	101
S5	103.5	0	111.5	89.5	26.5	0	0	28	69	101
S6	103.5	0	71	89.5	0	0	0	0	69	101
S7	103.5	0	71	89.5	0	0	0	0	41	101
S8	100	0	71	75	0	0	0	0	41	101
S9	68.5	0	71	75	0	0	0	0	41	101
S10	68.5	0	71	75	0	0	0	0	41	91.5
S11	68.5	0	71	75	0	0	0	0	41	91.5
S12	68.5	0	46.5	75	0	0	0	0	41	91.5
S13	68.5	0	46.5	75	0	0	0	0	41	91.5
S14	68.5	0	27.5	75	0	0	0	0	41	91.5
S15	68.5	0	10	0	0	0	0	0	41	91.5
S16	11	0	0	0	0	0	0	0	41	91.5
S17	11	0	0	0	0	0	0	0	36.5	33
S18	0	0	0	0	0	0	0	0	36.5	33
S19	0	0	0	0	0	0	0	0	36.5	33
S20	0	0	0	0	0	0	0	0	36.5	0
S21	0	0	0	0	0	0	0	0	0	0

Forrás: Saját szerkesztés

Az f_{10} következő lépcsőjét három auditjelentéssel lehet elérni, melyet az f_{10} követelménye szerint csökkenő sorrendbe rendezve, a hibrid GBM algoritmust futtatva, az alábbi új rekord (o22) szűrhető be az adattáblába (29. táblázat), melynek utolsó oszlopa a jelenleg meghatározott iránypreferenciák mentén történő elmozdulás hatását jelenti.

29. táblázat: Az újonnan beszűrt rekord értékei és az eddigi legjobb rekordtól történő elmozdulás minősítése

Attribútumok	O22 leíró adatai	Elmozdulás hatása
f_1	10.8	Ideális
f_2	8.8	Semleges
f_3	49.5	Ideális
f_4	26.0	Ideális
f_5	51.0	Semleges
f_6	13.0	Semleges
f_7	72.0	Semleges
f_8	12.0	Semleges
f_9	3.1	Ideális
f_{10}	1.9	Ideális
Y	619.25	Ideális

Forrás: Saját szerkesztés

Az Y, azaz F1-Pont értéke ezen tanuló halmaz felhasználásával kerekítve 0.62-re nőtt, így már az első iteráció hatására, tudatos javítással, sikerült az eddigieknél kedvezőbb performancia értéket elérni, azaz a tanulási halmaz felhasználása elmozdult az optimális felé.

Az új rekord adattáblába való helyezése után az iteráció ismétlődik. Az irány-preferencia konfirmálása (vagyis ezen módszertani logikában az irányok optimális finomhangolásának lehetősége) és rangsorolás után az 14. számú mellékletben közölt adattáblára történik a termelési függvény illesztése, ahol a kiindulópont az új o22 rekord (30. táblázat).

A táblából kiolvasható, hogy az új rekord beillesztésével immáron az f_9 változó bizonyul temporálisan hatástalannak a célváltozóra nézve, valamint a performancia növelésének érdekében az f_8 (jogszabályi vizsgálatok száma) független változót szükséges változtatni, melyhez elegendő egyetlen auditjelentés eltávolítása a tanulóhalmazból. Az is megállapítható, hogy a legnagyobb potenciált (227 pont) az f_1 jelenti, ezért olyan auditjelentést érdemes elvonni, mely segít előkelőbb helyezést elérni az f_1 attribútumban, azaz, ha létezik ilyen auditjelentés, mely képes több attribútum elvárásainak eleget tenni, a kereső eljárásnak kötelező jelleggel azt kell kiválasztania. A legelőkelőbb helyen lévő súlyszámok összege 1160, tehát az új rekord beillesztésével úgy tűnik, hogy még nagyobb tér lehet a növekedésre. Az amplitúdó növekedés ismert elméleti maximum esetén, azonban, nem jelent további hermeneutikai erőteret.

30. táblázat: A populáció súlyszámai a második iterációban

Lépcsők	f_1	f_2	f_3	f_4	f_5	f_6	f_7	f_8	f_9	f_{10}
S1	232.5	34.5	209	157.5	8	338	32	135	0	13.5
S2	232.5	34.5	209	157.5	8	229.5	32	135	0	13.5
S3	177	34.5	209	157.5	1	0	32	131.5	0	0
S4	177	34.5	203.5	145	1	0	32	131.5	0	0
S5	177	17	171	145	1	0	9.5	131.5	0	0
S6	177	17	171	145	1	0	9.5	0	0	0
S7	177	17	154	115.5	1	0	9.5	0	0	0
S8	173.5	11.5	154	115.5	1	0	9.5	0	0	0
S9	100	0	146	31	1	0	9.5	0	0	0
S10	100	0	146	31	1	0	9.5	0	0	0
S11	100	0	105.5	31	1	0	9.5	0	0	0
S12	100	0	105.5	31	1	0	9.5	0	0	0
S13	100	0	104	31	1	0	9.5	0	0	0
S14	68.5	0	71	31	1	0	9.5	0	0	0
S15	5.5	0	65	31	1	0	9.5	0	0	0
S16	5.5	0	65	0	0	0	9.5	0	0	0
S17	1.5	0	65	0	0	0	0	0	0	0
S18	1.5	0	65	0	0	0	0	0	0	0
S19	0	0	65	0	0	0	0	0	0	0
S20	0	0	0	0	0	0	0	0	0	0
S21	0	0	0	0	0	0	0	0	0	0
S22	0	0	0	0	0	0	0	0	0	0

Forrás: Saját szerkesztés

A feltételeket kielégítő auditjelentés elvétele után az új tanulóhalmazon futtatva a kiválasztott gépi tanuló algoritmust az F1-Pont értéke 0.62-ről 0.63-ra módosul (öt tizedesjegyig kerekítve a módosulás 0.61925-ről 0.62762-re javul), tehát ismételten a második iterációban is sikeres volt a tanuló halmaz optimum közelítése, azaz a genetikai potenciál felé történő elmozdulása. Mivel elvileg többféleképpen is lehet új rekordot generálni elméletileg, így az is vizsgálható, vajon mely szcenárió milyen pontossággal követi az Y-ra gyakorolt becsült hatást a tényleges Y-nal összevetve. Ez a tény egy optimalizálás-hatékonysági kérdést is felvet, azaz mi a minimális keresési lépésszámhoz szükséges új-rekord-generálási stratégia – vagyis van-e genetikai potenciálja a következő rekordot kijelölő lépésnek, illetve, igaz-e, hogy ez a stratégia vezet a legkevesebb lépéshez a keresés során, mely megoldása nem része jelen dolgozatnak.

Közbenső iterációkat tartalmi okok miatt nem részletezve, az Y értéke 0.64-ben (öt tizedesjegyig kerekítve: 0.63673) látszik véglegesedni a 11. iterációban, ahol az első helyen szereplő súlyszámok összege már csak 0.71, messze az elméleti küszöbérték alatt immár. Következésképp, elméletileg még mindig megkereshető legalább egy jobb eredmény, azonban az f_1 esetén optimum figyelhető meg (a korrelációs együtttható értéke közelít a 0-hoz), ezért az algoritmus megszakítása következett be, melyet automatikusan is képes a rendszer közölni a trendfüggvény monitorozásán keresztül. Egyértelműen megállapítható, hogy az f_1 attribútumnak hatása van a célváltozóra, azonban a kettő lineárisan nem korrelál.

A 11. iterációhoz tartozó számítási mátrixokat a 15. és 16. számú mellékletek szolgáltatják. Megfigyelhető, hogy 4 auditjelentést kellett véglegesen eltávolítani a tanulóhalmazból, hasonlóan, ahogy ez a második iterációban is történt. Ez azt jelenti, hogy volt olyan pont (10. iteráció), ahol az algoritmus egy iterációban az eddig kivett auditjelentéseket ismételten felhasználta már egy korábbi kivétel után, mert az előre történő elmozdulás kizárólag ily módon volt lehetséges. Amennyiben, már az első lépésben a kereső eljárás addicionális mintát igényelt volna feldolgozásra, a rendszer kilépett volna az iterációból és azzal a konklúzióval tért volna vissza, hogy jelen attribútumok és alappopuláció mellett optimálisnak tűnik a tanulás. Ebben az esetben az attribútumok és alappopuláció inicializálásának átgondolása egy lehetséges út, ritka esetben az is előfordulhat, hogy a rendelkezésre álló adathalmaz nem tartalmaz zajokat, bármely elvonásával a teljesítménymutató kedvezőtlenebbé válik.

Keresés befejezése

A kereső eljárás az alábbi táblázatokban (31., 32. és 33. táblázat) közöltek szerint javította a hibrid ABM, GBM és NN algoritmusok teljesítményét, ahol az ABM esetében már az első iterációban megtörtént a maximum megtalálása, míg az NN módszernél ez a második iterációban következett be, minimális javulást hozva. Kiemelendő, hogy a heurisztika nem feltétlen a lehető legjobb megoldást találja meg, mivel a keresés minőségét alapvetően befolyásolja az alappopuláció inicializálása.

31. táblázat: A hibrid és a kereső eljárással javított ABM performancia metrikái

Performancia mutatók	Hibrid ABM	Javított Hibrid ABM	Változás mértéke (%)	Cél	Változás hatása (javulás/romlás)
<i>Igaz pozitívak száma (db)</i>	89	102	+14.60%	Növelés	Javulás
<i>Igaz negatívak száma (db)</i>	944	940	-0.42%	Növelés	Romlás
<i>Hamis pozitívak száma (db)</i>	19	23	+21.05%	Csökkentés	Romlás
<i>Hamis negatívak száma (db)</i>	70	57	-18.57%	Csökkentés	Javulás
<i>Pontosság</i>	0.92	0.93	+1.09%	Növelés	Javulás
<i>Precizitás</i>	0.82	0.82	0%	Növelés	Semleges
<i>Fedés</i>	0.56	0.64	+14.29%	Növelés	Javulás
<i>F1-Pont</i>	0.67	0.72	+7.46%	Növelés	Javulás
<i>Variancia</i>	0.06	0.06	0%	Csökkentés	Semleges
<i>AUROC</i>	0.90	0.86	-4.44%	Növelés	Romlás
<i>AUPRC</i>	0.70	0.71	+1.43%	Növelés	Javulás

Forrás: Saját szerkesztés

32. táblázat: A hibrid és a kereső eljárással javított GBM performancia metrikái

Performancia mutatók	Hibrid GBM	Javított Hibrid GBM	Változás mértéke (%)	Cél	Változás hatása (javulás/romlás)
<i>Igaz pozitívak száma (db)</i>	62	78	+25.81%	Növelés	Javulás
<i>Igaz negatívak száma (db)</i>	958	955	-0.31%	Növelés	Romlás
<i>Hamis pozitívak száma (db)</i>	5	8	+60.00%	Csökkentés	Romlás
<i>Hamis negatívak száma (db)</i>	97	81	-16.49%	Csökkentés	Javulás
<i>Pontosság</i>	0.91	0.92	+1.10%	Növelés	Javulás
<i>Precizitás</i>	0.93	0.91	-2.15%	Növelés	Romlás
<i>Fedés</i>	0.39	0.49	+25.64%	Növelés	Javulás
<i>F1-Pont</i>	0.55	0.64	+16.36%	Növelés	Javulás
<i>Variancia</i>	0.03	0.03	0%	Csökkentés	Semleges
<i>AUROC</i>	0.85	0.85	0%	Növelés	Semleges
<i>AUPRC</i>	0.71	0.68	-4.22%	Növelés	Romlás

Forrás: Saját szerkesztés

33. táblázat: A hibrid és a kereső eljárással javított NN performancia metrikái

Performancia mutatók	Hibrid NN	Javított Hibrid NN	Változás mértéke (%)	Cél	Változás hatása (javulás/romlás)
<i>Igaz pozitívak száma (db)</i>	76	80	+5.26%	Növelés	Javulás
<i>Igaz negatívak száma (db)</i>	926	926	0%	Növelés	Semleges
<i>Hamis pozitívak száma (db)</i>	37	37	0%	Csökkentés	Semleges
<i>Hamis negatívak száma (db)</i>	83	79	-4.82%	Csökkentés	Javulás
<i>Pontosság</i>	0.89	0.90	+1.12%	Növelés	Javulás
<i>Precizitás</i>	0.68	0.68	0%	Növelés	Semleges
<i>Fedés</i>	0.47	0.50	+6.38%	Növelés	Javulás
<i>F1-Pont</i>	0.56	0.58	+3.57%	Növelés	Javulás
<i>Variancia</i>	0.11	0.11	0%	Csökkentés	Semleges
<i>AUROC</i>	0.85	0.81	-4.71%	Növelés	Romlás
<i>AUPRC</i>	0.62	0.62	0%	Növelés	Semleges

Forrás: Saját szerkesztés

A táblázatok szemléltetik, hogy az algoritmusok egyes mutatói romlottak az F1-Pont javítása érdekében, mivel a kereső eljárás az F1-Pont ideálisabb értékének letapogatására szolgált a többi metrika figyelmen kívül hagyásával. Kivétel nélkül ez az igaz pozitív (és így a hamis negatív) találatok arányának javításával volt elérhető. Az ABM és GBM alapú rendszerek igaz negatív és hamis pozitív találatainak száma csökkent, míg az NN esetében nem változott. Az AUROC mutatók a GBM kivételével romlottak, míg a varianciára az eljárás nem volt hatással, azaz ugyanolyan ütemben javult a tanulás pontossága, mint a tesztelése. Összességében kijelenthető, hogy a kereső eljárás sikerrel járt, mind a három algoritmus esetén javult az F1-Pont a tanulóhalmaz optimumhoz közelebb felhasználásával.

4.3.3. Objektív modelljóság-becslés a javított modelleken anti-diszkriminatív eljárással

Hasonlóan a 4.2.8. alfejezetben is közöltekkel, a javított algoritmusokra is végezzük el az anti-diszkriminatív elemzést az objektív megerősítésre, ahol a 25. táblázatban bemutatott modellek közé illesszük be az új javított modelleket (GP előtaggal ellátva), majd futtassuk újra a hasonlóságelemzést, ahol a norma értéke ismételtén 1000. Az eredményeket a 34. táblázat szemlélteti, kijelölve az az összehasonlítás alapját képező alkalmazásokat (31., 32. és 33. táblázatban közölt modellek).

34. táblázat: Modelljóság-becslés anti-diszkriminatív eljárással

Modellek	Pontosság	Precizitás	Fedés	F1-Pont	Variancia	AUROC	AUPRC	Y₀
<i>Egyszerű ABM - TA</i>	0.89	0.66	0.50	0.57	0.08	0.85	0.58	991.60
<i>Egyszerű GBM - TA</i>	0.89	0.83	0.29	0.43	0.04	0.85	0.61	965.50
<i>Egyszerű NN - TA</i>	0.88	0.61	0.47	0.53	0.11	0.82	0.56	952.50
<i>Hibrid ABM - TA</i>	0.92	0.82	0.56	0.67	0.06	0.90	0.70	1031.70
<i>Hibrid GBM - TA</i>	0.91	0.93	0.39	0.55	0.03	0.85	0.71	1022.70
<i>Hibrid NN - TA</i>	0.89	0.68	0.47	0.56	0.11	0.85	0.62	974.60
<i>Egyszerű ABM - KKIV</i>	0.88	0.67	0.60	0.63	0.09	0.88	0.71	1002.70
<i>Egyszerű GBM - KKIV</i>	0.89	0.73	0.59	0.65	0.03	0.90	0.75	1030.70
<i>Egyszerű NN - KKIV</i>	0.89	0.75	0.65	0.63	0.08	0.91	0.74	1028.70
<i>Hibrid ABM - KKIV</i>	0.94	0.90	0.73	0.80	0.06	0.91	0.84	1061.80
<i>Hibrid GBM - KKIV</i>	0.93	0.94	0.64	0.76	0.02	0.93	0.88	1065.80
<i>Hibrid NN - KKIV</i>	0.91	0.80	0.59	0.68	0.07	0.90	0.78	1033.70
<i>Egyszerű ABM - PSZ</i>	0.88	0.57	0.45	0.50	0.10	0.78	0.49	947.00
<i>Egyszerű GBM - PSZ</i>	0.90	0.78	0.32	0.46	0.05	0.79	0.52	977.10
<i>Egyszerű NN - PSZ</i>	0.87	0.51	0.35	0.42	0.09	0.75	0.38	940.50
<i>Hibrid ABM - PSZ</i>	0.90	0.73	0.45	0.55	0.08	0.80	0.55	982.60
<i>Hibrid GBM - PSZ</i>	0.91	0.89	0.38	0.54	0.04	0.82	0.63	1003.70
<i>Hibrid NN - PSZ</i>	0.87	0.54	0.32	0.40	0.09	0.76	0.40	941.00
<i>GP ABM – TA</i>	0.93	0.82	0.64	0.72	0.06	0.86	0.71	1039.80
<i>GP GBM – TA</i>	0.92	0.91	0.48	0.64	0.03	0.85	0.68	1030.70
<i>GP NN - TA</i>	0.90	0.68	0.50	0.58	0.11	0.81	0.62	975.60

Forrás: Saját szerkesztés

Összehasonlítva az Y_0 által szolgáltatott értékeket, belátható, hogy objektíven is kedvezőbb eredményt lehet elkönyvelni összesítve, a javított hibrid modellek mindegyike ideálisabb, mint a teljes adatvagyonon tanult hibrid modelleké, tehát a kereső eljárás sikerrel járt.

4.3.4. Az alfejezet összefoglalása

Belátható, hogy az ismertetett algoritmus kényszerűen képes megtalálni az ideális célváltozóhoz vezető utat, amennyiben adott iterációban rendelkezésre álló tanulóhalmaz attribútumai és a célváltozó között összefüggés tapasztalható, mivel meghatározott mértékű elmozdulás esetén ismertté válik a célváltozó értékek különbsége, mely adatokra az eljárás optimális választ szolgáltat, tehát az eljárás nélkül a keresés nem kivitelezhető. Sőt, az attribútumok együttes viselkedése és hatásai is mérhetővé válnak az iterációk folyamán, azaz minden egyes iteráció egy kvázi érzékenységvizsgálat a célváltozó értékére. Minden ismételt rögzített rekord újabb információval szolgál a termelési függvények részére, így véletlen attribútumok esetén azok hatástalansága beigazolódik.

Mint a legtöbb kereső eljárás, a bemutatott algoritmus sem képes a tökéletes célváltozó meghatározására, mivel ahhoz a teljes kombinatorikai teret lenne szükséges bejárni, azonban már egy pozitív véletlenül választott keresési szcenárió által is bizonyított, hogy:

- A keresési eljárás képes meghatározni, hogy milyen attribútumok módosítása szükséges az ideális célváltozó irányába történő elmozduláshoz;
- A kereső eljárás véletlen mutáció és az egyedek keresztezése, valamint köztes populációk definiálása nélkül is képes ideálisabb irányt megjelölni;

A leírtak alapján kijelenthető, hogy az alfejezet elején ismertetett hipotézis igazolható.

4.4. Modell-preferencia levezetése klasszikus tesztelési eljárások nélkül

Az alfejezet célja a tesztelés nélküli modell-preferencia matematikai levezetése, az alábbi hipotézis igazolása:

H3: A mesterséges intelligenciával ellátott döntéstámogató rendszerek teljesítményalapon a gépi tanuló alkalmazások klasszikus tesztelési eljárásai nélkül is rangsorolhatók, a predikciók, mint generált gyanúforrások leíró tulajdonságainak érték-irány levezetésével és az ezen adatokat feldolgozó matematikai apparátussal, mely automatizáltan képes a preferált modellek objektív meghatározására.

A gépi tanuló algoritmusok szakirodalma számos modelljóságot leíró mérőszámot határoz meg (többek között a dolgozatban is használt és ismertetett: Pontosság, Precizitás, Fedés, F1-Pont, stb.) a tanuló eljárások teljesítményeinek minősítésére, azonban mindegyikben közös, hogy a jósgmérés egy előzetesen a tanulásba be nem vont, azaz a tanulás érdekében nem hasznosítható tesztelési/validációs halmazon történik. Ennek előnye, hogy az eredmények értékelésére független adathalmaz kerül felhasználásra, ebből adódóan minimalizálható a túlilleszkedés/túltanulás jelensége, mely köszönhető annak, hogy az alkalmazott modell a tanuló adatbázisból megtanulta a felesleges zajokat, vagy éppen a tanulóminta közötti összefüggések leképezésére volt csak alkalmas, nem tudott kellően általánosítható tudást matematikai formába önteni. A tesztelési adathalmaz elkülönítésének létezik egy megkerülhetetlen hátránya, mégpedig, hogy értékes információk vagyontól (további tanulási mintáktól) fosztjuk meg a modelleket, azaz olyan mintát teszünk félre tesztelési célokra, melyből a modell még tanulni tudott volna. A szűkített elemszámmal ellátott tanulási halmazon kialakuló teszteredmények általánosságban preferáltabb eredményt mutatnak a tanuló adathalmazon végzett tesztelési eredményekhez képest, mivel a rendelkezésre álló értékes információk vagyont nem képes a tanulási folyamatba beépülni. Olyan területeken, ahol nehézkes a többletadat beszerzése vagy létrehozása, ez a megközelítés a modell megbízhatóságára is hatással lehet, s ilyen a kutatás kivitelezéséhez rendelkezésre álló adatvagyont is. Idősoros adatok esetén a tesztadatok kijelölése eleve zavaros: ha a távoli múltból tanulunk és a közelmúltot tesztelünk, akkor miért várható egyáltalán, hogy a közeljövőre a modell elég jó lesz, ha éppen az utolsó ismert történések nem kerültek megtanulásra. Emellett, ha addig tesztelünk, amíg a tanulás és a teszt összehangját megfelelően tarthatjuk, akkor lényegében (nagyon bonyolult módon), de a tesztadatok is megtanulásra kerülnek. Vagyis a tesztelés-alapú validáció egy fajta önámítás, vagyis a tesztelés nélküli validáció eleve zsákutcs elvárás. A kérdés csak az, miként lehet a tesztelés hagyományos lépéseit a minden adatot egyszerre „tanuló” rendszerben racionálisan kezelni?

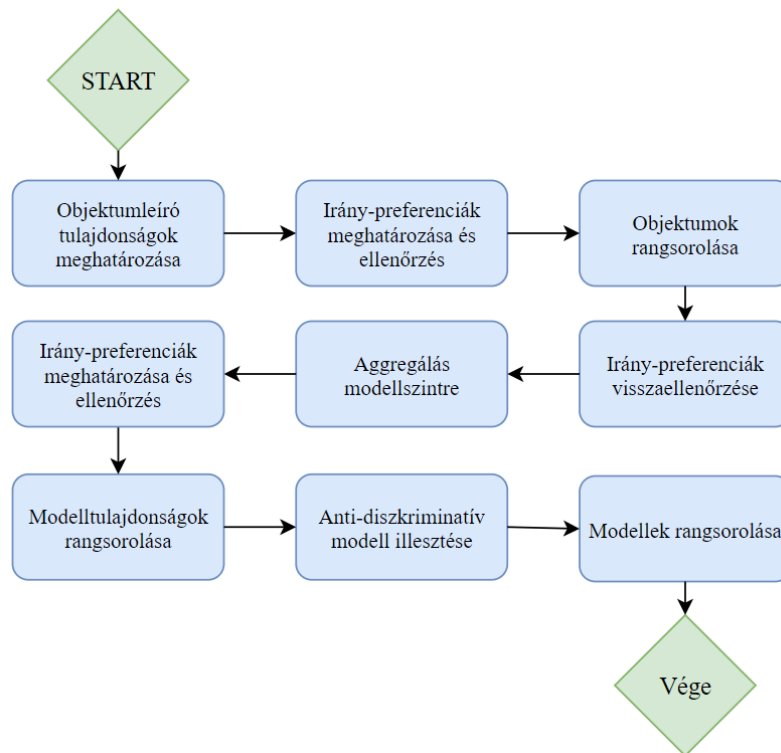
A tesztelés nélküli jósgmérés feltett szándéka, azaz a vonatkozó hipotézis célkitűzése, hogy bizonyítsa, hogy a fentebb említett probléma áthidalható azáltal, hogy a rendszer révén generált gyanús objektumok leíró tulajdonságaiból létrehozunk egy olyan adathalmazt, melyben több dimenzió mentén meghatározható, mely rendszerválaszokat preferáljuk (preferálja pl. a mindenkor döntéshozó, vagy a döntéstámogató rendszer maga) egy másiknál jobban. Tehát, melyek azok a jellemzők egy-egy rendszerválasz esetén, amelyek magasabb bizonyosságot (vö. konzisztenciát (PITLIK et al. 2017)) nyújthatnak a döntéstámogató rendszer felhasználójának. Egy olyan alkalmazás, mely nagy mennyiségű gyanút generál különböző kontrollokról, különböző kontrollterületekről, vélhetően elbizonytalanítja a döntéshozót és magasabb hamis pozitív eredményhez vezethet. Az alkalmazás, mely egységesen és homogén módon képes egy irányba mutatni (pl. egy kontrollt jelöl meg markánsan, mint gyanús eset, vagy hasonló struktúrájú/működésű kontrollok halmazát javasolja felülvizsgálatra) annak predikcióját vélhetően

könnyebben lehet elfogadtatni, mivel a rendszer minden egyes részegysége más utak bejárásával is hasonló konklúzióra jut. A meghatározandó leíró tulajdonságok, ennél fogva, a sikeres és sikertelen modellezés minél szélesebb körű aspektusait kell, hogy definiálják modelljóságot leíró mérőszámok formájában. Mivel ezen aspektusok száma végtelen, így a klasszikus egytényezős optimalizálást kényszerűen fel kell, hogy váltsa egy többdimenziós érték-fogalom a Hartman-elvet finomhangolandó (PITLIK et al. 2020c). Természetesen a leírtak nem jelentik azt automatikusan, hogy minden esetben pontatlanabb az az alkalmazás, mely bizonytalanabb és több alternatívát is felajánl, csupán azt, hogy elvárjuk egy döntéstámogató rendszertől, hogy rendelkezzen önerősítő mechanizmusokkal, mely különböző egymástól független eljárásokat is alkalmazva következetes döntéssel támogat.

A hipotézis akkor tekinthető bizonyítottnak, ha a fejlesztett modellek által generált gyanús outputok, azaz kontrollok halmazát irány-preferenciával ellátó leíró tulajdonságok automatizált elemzése és értékelése révén, sikerül olyan rangsort felállítani a modellek között, mely képes visszatükrözni a döntéshozói elvárást (pl. Turing-teszt), következésképp, az a modell kerüljön kiválasztásra, mely képes holisztikusan helyesebben megítélni egy gyanú létjogosultságát. Kiemelendő, hogy a mindenkori cél az objektivitás kikényszerítése, azaz a döntéstámogató rendszer a modell-értékelésig nem engedheti meg a döntéshozó beavatkozását, hacsak nem a rendszer döntésképtelen állapotba jut (vagyis „nem tudom” rendszerválasszal tér vissza). A döntésképtelenség nem csak azt jelentheti, hogy a rendszer nem tud válaszolni a kérdésekre, hanem azt, hogy az általa megoldandó feladat végeredménye az alkalmazott matematikai apparátusokon keresztül egy olyan ítélet, mely a felhasználó preferenciájától függ, azaz objektíven, racionális indokokkal alátámasztva egy robot szemszögéből a robosztus alternatívák az eredményhirdetésnél nem különböznek. Továbbá, az irány-preferencia a leíró tulajdonságok tekintetében az objektíven elvárt rendszerválaszok jóságát kell, hogy tükrözzék, tehát racionálisnak kell lenniük, azonban, a többdimenziós jóságfogalom eredménye, lényegében a konszenzuskötés matematikáját követve egyetlen egydimenziós döntéshozatal eredményével sem kell, hogy kényszerűen azonos legyen (PITLIK et al. 2020c).

4.4.1. A javasolt eljárás

A modell-preferenciák levezetésére javasolt eljárás folyamatábráját a 30. ábra vizuálisan szemlélteti.



30. ábra: Modell-preferenciák levezetésére javasolt eljárás folyamatábrája

Forrás: Saját szerkesztés

A modell-preferenciák levezetésére javasolt eljárás az alábbiakban leírt lépésekben képes az ideális modell megtalálására:

1. **Objektumleíró tulajdonságok meghatározása:** A modellekhez, mint objektumokhoz kapcsolódó rendszerválaszokat leíró attribútumok meghatározása, melyek a gyanúgenerálásra adott válaszokat minősítik tetszőlegesen sok (racionális) szempontból. Együttes feldolgozásuk képes a rendszerválaszokban rejlő ellentmondások felfedezésére, ezáltal az adott modell illogikus mechanizmusát a későbbiekben büntetőponttal látja el (pl. modell aggregálásakor). Az attribútumok létrehozásával előáll egy objektumszintű adattábla.
2. **Írány-preferenciák meghatározása és ellenőrzése:** Az objektumleíró tulajdonságok irány-preferenciát szükséges meghatározni statisztikai módszerekkel, mely hozzájárul az összefüggések (attribútumok és teljesítmény) helyesnek vélt irányaihoz, ahol a Hartman-i értékfogalom meghaladása lehetővé teheti az irányok optimalizálását is (PITLIK et al. 2020c).
3. **Objektumok rangsorolása:** Az irány-preferenciák alapján állítsuk fel objektumszinten a rangsort, ahol az ideálisnak vélt irány alapján meghatározható, hogy mely objektum teljesít adott attribútum elvárásai szerint jobban/rosszabbul a többi objektumhoz képest.
4. **Írány-preferenciák visszaellenőrzése:** A rangsorolás után szükséges az irány-preferenciák ismételt ellenőrzése, azaz azok érvényességének vizsgálata inverz nézetben, mely szolgáltathatja azt az eredményt, hogy az eddig helyesnek vélt irány-preferencia alapján nem teljesülnek a függvény-szimmetria által támasztott követelmények (nem érvényes adott

objektum), ahol a fekete-fehér validitás-vizsgálat speciális fuzzy-alakzattá is formálható (PITLIK L. – PITLIK M. 2021b)

5. **Aggregálás modellszintre:** Az adattáblát aggregáljuk objektumszintről modellszintre egyváltozós statisztikai eljárások segítségével pl. átlag, maximum, minimum, szórás stb., ahol modellszintű metrikák állnak rendelkezésre a modellek tulajdonságainak leírására.
6. **Irány-preferenciák meghatározása és ellenőrzése:** Határozzuk meg a modellszintű adattábla attribútumainak irány-preferenciáit, majd ellenőrizzük érvényességüket.
7. **Modelltulajdonságok rangsorolása:** Rangsoroljuk a modelleket az előző pontban meghatározott irány-preferenciák alapján.
8. **Anti-diszkriminatív modell illesztése:** A modellszintű adattáblára illesszünk anti-diszkriminatív lépcsős függvényt, melyben a célváltozó konstans, kikényszerítve a normától való eltérés beazonosíthatóságát (a normától való eltérés jelentheti azt, hogy egy adott modell meghaladja a többi modell teljesítőképességét, azaz a leíró tulajdonságok adott modell esetén a legideálisabbak).
9. **Modellek rangsorolása:** Vizsgáljuk meg az anti-diszkriminatív lépcsős függvény becsléseit, hirdessük ki győztesnek a legmagasabb értéket. Amennyiben teljesül a „mindenki másképp egyforma elv”, töröljük a már felhasznált attribútumokat és ismételten illesszünk anti-diszkriminatív lépcsős függvényt az adattáblára. Másik megoldás, hogy a döntéshozó elfogadja, hogy a modellek lényegében nem különböznek, mert mindegyik modell „győztes” pl. egy adott „versenyszámban”.

Az algoritmus pszeudokódja megtalálható a 11. mellékletben.

4.4.2. A modellezés folyamata, alkalmazott gépi tanuló eljárások

A kutatásban a 4.2.5 alfejezetben ismertetett felügyelet nélküli algoritmusokat alkalmaztam, melyekhez további alternatívákat fejlesztettem, rendszersúlyozással és anélkül, valamint megállapítássúlyozással és anélkül, így 12 modell kerül felhasználásra az optimalizálás igényét elsőként nem preferálva, tehát a cél nem a problémára illeszkedő legjobb megoldás elkészítése volt, hanem az, hogy bizonyítsam, hogy kvázi véletlen modellek között is megtalálható a legideálisabb. A feladat megoldására a már elkészített felügyelt gépi tanuló módszerek bizonyulnak a legkiválóbbnak szemléltetésül, mivel képesek kilistázni, akár redundánsan is, kimenetként gyanúmomentumokat, ezért az önerősítő mechanizmusok megléte/hiánya egyszerűbben validálható a dolgozatban alkalmazott felügyelt módszerekkel ellentétben.

Legyen $X \rightarrow Y$ egy olyan matematikai leképezés, ahol X jelenti az input, Y az output teret. Fejezze ki $x^{(i)}$ az X i -edik objektumát, $f^{(i)}$ az X i -edik attribútumát. Jelölje p azt a vektort, mely tartalmazza $\forall x \in X$ -re a $d(x^{(i)}, x^{(j)})$ -ből képzett vektornak $k \in \mathbb{N}$ legnagyobb elemét, ahol d a kiválasztott távolság/kapcsolat nagyságát meghatározó függvény, k pedig tetszőlegesen kiválasztott. Így, k értéke szabályozni képes a generált gyanúmomentumok sokaságát, tehát csökkentésével kikényszeríthető a kevesebb rendszerválasz, mely az igaz pozitív találatokat is csökkentheti. Növelésével a hamis pozitív arányt képes növelni, ezért a k megválasztása ebben a pillanatban a

döntéshozói preferenciát tükrözi, figyelembe véve az elérhető erőforrásokat a gyanús esetek kivizsgálására. A k értékének megtalálása, mindazonáltal automatizálható.

A q vektor tartalmazza a p elemeihez tartozó x objektumait. Tehát a q és p vektorok egy adott kiválasztott auditjelentés (objektum) esetén tartalmazza a top k leghasonlóbb elemeit és azok számított távolságait. A Q mátrix tartalmazza az f címkéit, tehát azon kontrollokat, melyek a top k objektumok esetén a kiválasztott x és a q elemei között különbséget mutat, azaz gyanús kontroll. Amennyiben a kontroll nem volt az adott audit hatókörében, így azt jelölje -1, azaz a generált gyanú nem értelmezhető. Belátható, hogy a -1 egy megfelelő jelölés azon kontrollokra, melyek nem képezték hatókörét a vizsgálatnak, mivel a megállapítások számának pozitív egész számnak kell lennie. Az Y tehát egy olyan halmaz, amely tartalmazza a p , q vektorokat és a Q mátrixot és a generált gyanúk halmaza kizárólag p , q és Q alapján értelmezhető. A d a 4.2.5 alfejezetben ismertetett 3 különböző távolságot/kapcsolatot mértékét meghatározó függvény, azaz:

$$(1) \text{ EUC} = \text{Euklideszi távolság: } d(x^i, x^j) = \sqrt{\sum_{l=1}^n (x_l^i - x_l^j)^2}$$

$$(2) \text{ PEA} = \text{Pearson-féle korreláció: } d(x^i, x^j) = \frac{\sum_{l=1}^n (x_l^i - \bar{x}^i)(x_l^j - \bar{x}^j)}{\sqrt{\sum_{l=1}^n (x_l^i - \bar{x}^i)^2 \sum_{l=1}^n (x_l^j - \bar{x}^j)^2}}$$

$$(3) \text{ COS} = \text{Koszinusz hasonlóság: } d(x^i, x^j) = \frac{\sum_{l=1}^n x_l^i x_l^j}{\sqrt{\sum_{l=1}^n x_l^i{}^2} \sqrt{\sum_{l=1}^n x_l^j{}^2}}$$

A rendszerek súlyait számítsuk ki úgy, hogy az audit hatókörébe tartozó rendszerek számával osszuk el az összes megállapítást, így megkapjuk az egy rendszerre eső (relativizált) megállapítások számát. A rendszerek súlyozásának bevezetése racionális döntés lehet, mivel az auditjelentések összehasonlításakor a rendszer kikerüli a komplexebb objektumok (auditok) túldominálásának csapdáját, ha egy nagyobb számú informatikai infrastruktúrát vizsgáló auditot hasonlítunk össze egy kevésbé összetettebb audittal, ebből adódóan ugyanazon a skálán lehet értelmezni minden objektumot.

A megállapítások súlyának meghatározására az adott q vektor objektumaihoz tartozó kontrollokhöz (Q mátrix) rendeljünk egy Q' mátrixot, mely egy adott kontrollhoz tartozó megállapítások számának arányát fejezi ki a legnagyobb megállapításhoz képest, azaz az a kontroll lesz a legsúlyosabb (értéke 1), amelyikhez a legtöbb megállapítás tartozik. Ezáltal finomhangolhatóvá válik a gyanúgenerálás és felállítható egy preferencia lista, amit az auditor csökkenő sorrendben felülvizsgálhat a jelentés kiadása előtt.

A felhasznált adatbázis a kísérletben 53 auditjelentést tartalmaz (pénzügyi szektorban végzett könyvvizsgálathoz kapcsolódó informatikai vizsgálatokat), mely a legnagyobb homogén részhalmaza a teljes adatvagyonnak a kiélelt verseny fokozása végett, valamint 24 kontrollt, amely releváns hatóköre volt az auditoknak.

A k értéke az elemzésben 10, azaz a top 10 objektumot veszik a modellek alapul a gyanúgenerálásra, mely egy önhatalmú választás eredménye (hasonlóan a korábbi felügyelet nélküli modellezésnél a 4.2.5. alfejezetben). Mivel a 10 objektumba beletartozik az elemzésre kiválasztott auditjelentés is, ezért gyakorlatilag 9 objektum alapján fognak a modellek ítélni,

mely az összes objektum 16.98%-a, ezért várhatóan magas lesz a hamis pozitív találatok aránya. Összefoglalóan az alábbi, 35. táblázat szemlélteti a kísérletbe bevont modelleket a leírtak alapján.

35. táblázat: Felügyelet nélküli modellek leíró dimenzióinak összefoglaló táblázata

ID	Modell azonosító	Megállapítások súlyozása	Rendszerek súlyozása	Távolság/kapcsolat- metrika
1	R_NS_NRS_C	Nincs	Nincs	Koszinusz
2	R_NS_NRS_E	Nincs	Nincs	Euklideszi
3	R_NS_NRS_P	Nincs	Nincs	Pearson
4	R_NS_RS_C	Nincs	Súlyozott	Koszinusz
5	R_NS_RS_E	Nincs	Súlyozott	Euklideszi
6	R_NS_RS_P	Nincs	Súlyozott	Pearson
7	R_S_NRS_C	Súlyozott	Nincs	Koszinusz
8	R_S_NRS_E	Súlyozott	Nincs	Euklideszi
9	R_S_NRS_P	Súlyozott	Nincs	Pearson
10	R_S_RS_C	Súlyozott	Súlyozott	Koszinusz
11	R_S_RS_E	Súlyozott	Súlyozott	Euklideszi
12	R_S_RS_P	Súlyozott	Súlyozott	Pearson

Forrás: Saját szerkesztés

4.4.3. Feltételezések modelljóságra a megalkotott modellek alapján

A fejlesztett 12 modell a k érték magas megválasztása miatt feltételezhetően magas hamis pozitív aránnyal fog bírni (sokkal több a gyanúmomentum, mint a gyanús eset), melynek leleplezésére a modelljóságot kereső algoritmusnak alkalmasnak kell lennie a rendszerválaszok alapján. Mindamellet, a magas hamis pozitív arány döntéshozói szempontból preferált is lehet, ha ez az igaz pozitív arány növelésének javára történik. Egy olyan szervezet, mely rendelkezik elegendő erőforrással (idő, tőke és emberi) a hamis pozitív találatok feltárása alaposabb auditori munkát eredményezhet magasabb szintű bizonyossággal, ezért az audit vezetője dönthet ennek irányába.

Az erőforráskorlátokkal bíró auditor megelégedhet egy olyan megoldással is, mely alacsony számban képes gyanút megállapítani, de annak döntő többsége visszaigazolhatóan hibás audit eljárásnak volt köszönhető (alacsony igaz pozitív és alacsony hamis pozitív ráta). A rendszerek súlyozása lehetővé teszi, hogy egy adott auditra vonatkozóan ne egész szám legyen adott kontrollra eső megállapítások száma, ezért a pozitív eltérés is a hamis pozitívak számát növelheti.

A megállapítások súlyozása esetén szükséges meghozni egy olyan konszenzust, amely nyilatkozik az elfogadható határértékről, azaz mi az a súlymérőszám, amely felett elfogadjuk egy gyanús kontroll létét, alatta nem történik meg a kontroll felülvizsgálata. A modellezés során minden generált gyanú, amelynek súlya nem 0, a könnyebbség kedvéért figyelembevételre kerül, ezért feltehetően a megállapítások súlyozása további hamis pozitív eredményeket generálhat. Továbbá, mivel szűkített adathalmazról van szó szűkített attribútumokkal, így az Euklideszi távolságmétrikát alkalmazó algoritmusok potenciális győztesei lehetnek a versenynek, mivel két vektor közelebb eshet egymáshoz minimális eltérésekkel különböző attribútumokban, mint a Koszinusz hasonlóság és Pearson-féle korreláció esetén mért hasonlóságok/kapcsolatok, mely azonos attribútumokban mért egyezőségek/eltérések esetén preferáltak.

4.4.4. Objektumleíró tulajdonságok meghatározása

A visszaellenőrzésére (öniróniánkat megtartva), szükséges elkülöníteni egy teszthalmazt, hogy igazolva legyen, hogy objektumleíró tulajdonságok által is hasonló (illetve racionálisan eltérő) konklúzióra lehet jutni, mintha klasszikus tesztelési eljárásokkal bizonyosodtunk volna meg a legjobb modell kilétéről. Az 53 auditjelentésből különítsünk el 12-t (22.64%), mely halmazon mérni fogjuk a találati arányokat, és a halmazon végezzük el a leíró attribútumok kiértékelését. Ahhoz, hogy mérni lehessen, hogy valóban helyes konklúzióra jutottak-e a modellek, véletlenszerűen változtassunk meg 1-1 attribútumot úgy, hogy a valódi gyanú tényét elfedve azt mutatjuk a modellek felé, hogy adott kontroll esetén nem talált az auditor kivetni valót (szimulált környezet). Ez a visszaellenőrzés így valóban objektív, szemben a nyersadatok által és az utólag jelzett gyanúk beépített szubjektivitási terhével, tehát el fogjuk tudni dönteni, hogy a szándékosan manipulált attribútumot jelezték-e felénk a modellek, vagy sem. A modellek kiértékelésére a korábban is alkalmazott (3.3.1 alfejezet) mutatókat használom a Variancia kivételével, mely itt klasszikusan nem értelmezhető.

A modellek és a visszaellenőrzésre elkészített mutatók meghatározása után a következőkben hozzunk létre olyan, a modelljóságokat kifejező attribútumokat minden egyes kiválasztott modell és auditjelentés pároshoz, melyek irány-preferenciája alapján az egyes modellek rendszerválaszai külön-külön is rangsorolhatók és teljesül a „minél kisebb/nagyobb, annál konzisztensebb/jobbb” elv. A 17. számú melléklet a kutatáshoz megalkotott attribútumokat foglalja össze. Ahogy korábban is említésre került, az objektumleíró tulajdonságoknak racionálisnak, hermeneutikailag védhetőnek, egymáshoz képest is értelmezhetőnek kell lenniük ahhoz, hogy képesek legyenek együttes hatásmechanizmussal a „jó/jobbb/legjobb” kiértékelésére. Amennyiben sikerül elegendő számú attribútumot találni és azokat közösen feldolgozni, az egyes aggregált értékeknek korrelálniuk kell a jóságmutatókkal megengedve némi szórást, mivel nem szigorú fizikai adatmérésről lévén szó az adathalmaz, így tartalmazhat véletlen zajt. Ennek mértéke a Knuth-i elv alapján önmagában is egy optimalizációs feladat. Másodsorban, mivel az objektumleíró tulajdonságok száma végtelen bármilyen problémáról legyen is szó, ezért a tökéletes egyezés csak szómisztika. A megfelelő attribútumok alapvetően kontextus független jellegűek, így ezek fokozatosan bővülő katalógussá formálhatók és a mindenkori automatikus modellezés alapjaként felhasználhatók. A dolgozatnak nem fókusza a Knuth-i elv kiterjesztése újabb és újabb attribútumok automatizált fellelésére.

A 17. számú mellékletben ismertetett kvázi véletlenszerű mutatók képesek kezelni a „nincs megállapítás” tényét is, azaz elképzelhető olyan rendszerválasz, mely a döntéshozó számára azt jeleníti meg, hogy nem volt gyanús objektum. A T01 és T02 attribútumokkal tehát körbezártuk a rendszerválaszok mozgásterét, mivel az egyiktől azt várjuk el, hogy „beszéljen” azaz helyesebbnek ítéljük meg, ha generál gyanút, még ha az nem is igaz, míg a másiktól a minimális rendszerválasz a preferált.

A hibásan gyanúsított kontrollok tényét leleplezendő, több szempontból is mérésre kerülnek a homogenitás mutatók, melyek egységesen képesek nyilatkozni a többi auditjelentés alapján meggyanúsított kontrollokról, tehát kollektív „hazudás” szükséges, ahhoz, hogy egy rosszabb eredményt helyesnek ítéljünk meg, melynek valószínűségét a megválasztott k értéke tovább csökkenti. Evégett, a félrevezető modellek előbb-utóbb büntetőpontot kapnak, mely egyaránt alkalmas arra, hogy a rosszindulatból szándékosan manipulált auditjelentéseket kidobja magából a rendszer, így sikerült olyan attribútumokat kijelölni, mely akár a tagadás tagadásaként is szükségszerűen jutalmazza/bünteti az igazmondó és hazug modelleket.

A táblázat, továbbá tájékoztat arról is, hogy egy attribútum kontextus független vagy kontextusfüggő tulajdonsággal bír. A kontextus független tulajdonságokat a kiválasztott problémától függetlenül, szuverén módon is alkalmazni lehet, tehát nem követeli meg az adott szakterület átfogó ismeretét, annak hiányában is értelmezhetők. Ezen attribútumok minél pontosabb definiálása egy döntéstámogató rendszer számára kiemelten fontos, mivel ez hozzájárulhat egy általános problémamegoldó döntéstámogató rendszerhez, merész szavakkal élve az általános mesterséges intelligenciához, mivel hasonló megközelítéssel rendelkező feladatok elvégzése sablonosítható, a tudás transzferálhatóvá válik. A kontextusfüggő mutatók létrehozása és értékelése ezzel ellentétben szakterületi tudást igényel, a kijelölt feladatban ez azt jelenti, hogy az egymással szemantikailag (kontextusban) összefüggő kontrollok egymás hiányát/meglétét képesek erősíteni, ezért a problématerülettől függő mutatók elhanyagolása információvesztést okozhat. Legyen példa egy olyan szervezet, ahol nem megfelelően kivitelezett, vagy épp hiányzik a kriptográfiai szabályzat, ott feltehetően nincs vezetői nyomás és követelmény a kriptográfiai kontrollok betartására, így jó eséllyel az egyéb kriptográfiai kontrollok (pl. hálózati adatfolyam titkosítás, tanúsítványok kezelése, stb.) is hagy kivetni valót maga után.

4.4.5. Objektumleíró tulajdonságok irány-preferenciáinak ellenőrzése

Az algoritmusok futtatása és az objektum-attribútum adattábla felöltése után (ahol legyen T az objektumleíró tulajdonságok mátrixa), de bármilyen elemzés megkezdése előtt ellenőrizni szükséges az adattábla következetes meglétét, azaz valóban a megjelölt irány-preferenciák helyesen kerültek-e meghatározásra az ellentmondásosság kiküszöbölése érdekében, melyre pl. a korrelációs mátrix elemzése és/vagy a hasonlóságelemzés is alkalmas. A két módszer együttes alkalmazása segíthet az irány-preferenciák objektív meglétének megerősítésében vagy éppen elvetésében. A mesterségesen létrehozott objektumleíró tulajdonságok irány-preferenciái nem mindig egyértelműek, mert rendelkezhetnek nem monoton hatásmechanizmusokkal is, azaz egy adott attribútumnak optimuma is lehet (más, pl. periodikus hatásmechanizmusokat nem említve), ahogy ez a genetikai potenciált kereső algoritmus esetében is azonosításra került. Ez a gyakorlatban nem csak származtatott attribútumok esetén lehet releváns kérdés, hanem pl. ahogy ez később is bemutatásra kerül, a kidolgozott eljárásban a szórásmutatók irány-preferenciája sem mindig triviális. Itt és most nem arról kell nyilatkozni, hogy egy adott objektum adott attribútuma esetén biztosan a döntéshozó szempontjából van egy olyan meghatározott irány, mely általa előnyben részesített, hanem arról, hogy egy adott modell milyen utak bejárásával képes elérni „genetikai potenciálját”, azaz milyen módon válaszkapesebb egy modell a jót keresve. Ha a modellek, tételezzük fel, csak magas szórások révén hajlandóak növelni az igaz pozitív találatok számát, és az alacsony szórású modellek hibás fókusszal dolgoznak, azaz hibás feltételezéseket tesznek az összefüggésekre, akkor hiába az az elvárás, hogy a minimális szórás a legjobb, ha az a megkötés csak egy lokális optimum felé ágyaz meg az adott modellnek, mely a „no free lunch” elmélet alternatív értelmezésének tekinthető (WOLPERT 2020).

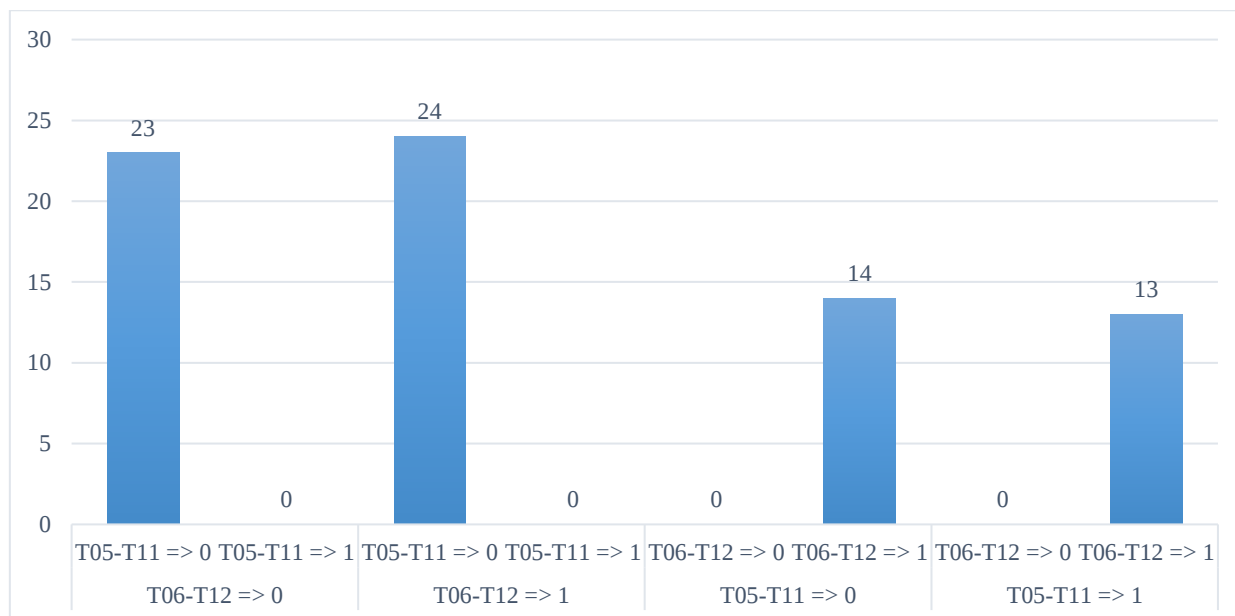
A példához visszatérve, a korrelációs együtthatókat elemezve, anomáliának nevezhetjük azt az értéket, ahol a hasonló irány-preferenciákat feltételező attribútumok között erős (legalább 0.7) (SAJTOS – MITEV 2007) negatív, míg a különböző irány-preferenciákat feltételező attribútumok között erős pozitív korreláció figyelhető meg. A mesterségesen létrehozott attribútumok közötti korrelációt a 18. számú melléklet szemlélteti. Irány-preferencia megjelölésre a számos szoftveres megoldásban is alkalmazott jelölésrendet alkalmaztam a táblázatban, mely 1, ha „minél kisebb, annál jobb” és 0, ha „minél nagyobb, annál jobb” érték. A felfedezett anomáliák kijelölésre kerültek, azaz azon pontok, ahol elvárnánk, hogy a különböző irány-preferenciával rendelkező

attribútumok korrelációs értékei egyenes/fordított irányt mutassanak. A T12 attribútum esetén markánsak az eltérések, és megfigyelhető az is, hogy bár egy esetben lépett fel erős korreláció, a többi attribútummal összehasonlítva az értékeket, a T06 úgy viselkedik, mintha a kezdetben meghatározott irány téves lenne. A T06 és T12 attribútumok már a 17. számú mellékletben ismertettek alapján is ki voltak jelölve további elemzésre, láthatóan az a sejtés, hogy az irány-preferencia tévesen került megítélésre. A T05 és T11 attribútumok alátámasztják a kezdeti irány-preferencia megjelölést.

Az irány-preferenciák kapcsán mindenkor figyelembe kell venni a kontextus-függőséget is: különösen a származtatott attribútumok esetén veszélyes olyan származtatás, ahol a minimum két nyers attribútum iránya/értelmezési tartományi gyengíti egymást (pl. rangsorszámozott értékek súlyozása, ahol a rangsorszám, mint olyan a „minél kisebb annál jobb” elvet fejezi ki, míg a súlyok általában a „minél nagyobb, annál jobb” elvet követik). A súlyozott sorszámátlag tehát irányíthatóság szempontjából egy nonszensz értelmezési teret generál.

Az irány-preferencia alátámasztására alkalmazzuk a hasonlóságelemzés által felkínált matematikai apparátust is. A hasonlóságelemzés képes a direkt és inverz nézetek párhuzamos jellegű elemzésére, azaz a függvény-szimmetria-sérülések leleplezésére (PITLIK et al., 2020c). Az objektumokat és azok leíró tulajdonságait tartalmazó táblázatban az objektumok rangsorolása által a hasonlóságelemzés idealitás-mutatót számol, mely alkalmas az adattábla vizsgálatára az eltérő inverz nézetekben a kérdéses attribútumok tekintetében, ahol a T05-T11 és T06-T12 párokat eltérő irány-preferenciával tanulmányozzuk, tehát összesen 4 különböző aspektusban. Mivel az irányok inverzéről van szó, ezért a hasonlóságelemzés alapján az egyes objektumoknak is illene szimmetrikusnak lenniük, ahol ez nem teljesül, ott a hasonlóságelemzés alapján érvénytelen objektum-értelmezhetőségről van szó. Ha egy adott nézetben csak/túlnyomóan érvénytelen hermeneutikájú objektum van, akkor az a vélelem, hogy adott álláspontból az az irány-preferencia is érvénytelen, vagy a nyersadatok véletlenszámok. Amennyiben az eltérő inverz aspektusokban hasonló az érvényes modellek száma (tehát az irány-preferencia továbbra sem dönthető el egyértelműen), ott az idealitás-mutatók átlaga nyújt segítséget, melyek statisztikai szignifikanciáját pl. varianciaelemzéssel szükséges alátámasztani. Az alábbi ábra szolgáltatja a hasonlóságelemzés által feltárt érvényes objektumok számát a 4 különböző inverz nézetben (31. ábra).

Kétségtől kijelenthető, hogy a hasonlóságelemzés is alátámasztja a T05-T11 „minél nagyobb, annál jobb” és a T06-T12 attribútumok „minél kisebb, annál jobb” irány-preferenciáját, mivel az eltérő inverz irányokban nem volt érvényes objektum. Az irányok inverzitása az összes attribútum esetén nagyobb értelmező erővel bír, mint az attribútumok korlátozott száma esetén, de egyetlen egy attribútum irány-inverzitása is hermeneutikai alapinformációként használható.



31. ábra: Felügyelet nélküli modellek irány-preferencia vizsgálata

Forrás: Saját szerkesztés

A két elemzést összevetve, objektíven is sikerült alátámasztást nyerni a preferált irányokról, melyet minden egyes attribútum esetén javasolt elvégezni, ahol kétség annak megfelelőségéről felmerül. Az értekezett problémában az irány-preferenciák matematikai úton történő levezetésének eredményét nem szabad örökérvényű igazságként elfogadni, ez csak azt jelenti, hogy ebben a kutatási példában, ezen adathalmazt és modelleket felhasználva a jó irányhoz vezető út vélhetően a meghatározott irány-preferenciák által érhető el. A módszer maga azonban kontextus független, vagyis a mindenkor valós helyzetben bármikor alkalmazható a szükséges paraméterek optimumának levezetése mellett. Az objektív megerősítés azért is fontos, mert az irány-preferencia erősen szubjektív lehet, azonban ez a kiértékelés is rámutat arra, hogy ez is automatizálható, tehát a döntéstámogató rendszer képes megmondani a modellezés szempontjából melyik út vezet(het) a jobbhöz.

Fontos visszautalni a H2 és H3 hipotézisek közötti kapcsolatra: a genetikai potenciál fogalmára alapozó innovatív keresésvezérlés még azt is megengedi, hogy iterációnként eltérő irány-preferenciák kerüljenek alkalmazásra ott, ahol ezt a kontextus maga nem zárja ki és az egy-egy rekordtöbbség korrelációs előjelváltásokat képes okozni.

A korrelációs mátrixot elemezve felmerülhet a multikollinearitás által potenciálisan okozható eredménytorzítás kockázata, azonban a modelljóság meghatározására alkalmazott hasonlóságelemzés ezt képes automatikusan kezelni és megjeleníteni az attribútumok súlyaiban (azaz kizárja/eltérően értékeli a redundáns attribútumhatásokat), ezért ez nem jelent fenyegetést, így minden attribútumot meghagyhatunk az elemzésben (PITLIK L. – PITLIK M. 2021a).

4.4.6. Modellek rangsorolása

A következőkben szükséges elkészíteni azt a mátrixot, melyet transzformálva az objektumleíró tulajdonságokból az irány-preferenciák alapján képes rangsorolni az egyes objektumokat azok attribútumai szerint, tehát:

$$T_r := \text{rank}(X, \xi)$$

ahol Tr jelenti az X mátrix rangsor szerinti transzformációját, ξ az irány-preferenciákat. A rangsoroláson kívül érdemes egyéb, iránnyal rendelkező adattranszformálást elvégezni, mely redundáns információt tartalmazhat, azonban egy új nézőpont szerint is képes kifejezni, melyik modell jobb egy bizonyos „versenyszámban” a másiknál, azaz nem kizárólag a nyers rangsor szerinti értékelés lesz az alapja a modell-idealitások meghatározásának. A percentilisalapú rangsorolás és a normalizált értékek kiszámítása alkalmas lehet erre a feladatra, ahol irány-preferenciának megfelelően kell direkt és inverz mutatókat létrehozni.

$$T_p := \text{percentrank}(X, \xi), \text{ ahol } \begin{cases} \frac{\max(X_j) - \text{rank}(X_{(i,j)}, 0)}{\max(X_j) - \min(X_j)}, \text{ ha } \xi = 0 \\ 1 - \frac{\max(X_j) - \text{rank}(X_{(i,j)}, 0)}{\max(X_j) - \min(X_j)}, \text{ ha } \xi = 1 \end{cases}$$

$$T_n := \text{norm}(X, \xi), \text{ ahol } \begin{cases} \frac{X_{(i,j)} - \min(X_j)}{\max(X_j) - \min(X_j)}, \text{ ha } \xi = 0 \\ 1 - \frac{X_{(i,j)} - \min(X_j)}{\max(X_j) - \min(X_j)}, \text{ ha } \xi = 1 \end{cases}$$

ahol i jelenti az X mátrix egy sorát, j jelenti az X mátrix egy oszlopát. Az így előállt adattábla, mely egyelőre modell-objektum szinten tartalmazza az értékeket, aggregálható modellszintre és modellenként meghatározható a jószágmutatók egyváltozós leírás statisztikái, mint pl. az átlag, maximum, minimum és szórás, melyet az 19. számú melléklet tartalmaz. Az átlag, maximum és minimum alapján triviális, hogy az irány-preferencia a „minél nagyobb, annál jobb” (illetve a rangsoroknál ez „minél kisebb, annál jobb”), ezzel szemben, a szórás esetén ez megkérdőjelezhető, ezért ellenőrizni szükséges a szórás-irányok megfelelőségét.

A modell jószág-szórások attribútumaira mindenképp szükség van, mivel nem mindegy, hogy egy adott modell következetesen minden tesztesetre hasonló rendszerválással tért vissza (még, ha hibás eredménnyel is), vagy magas szórás esetén a megtalált gyanú a véletlen műve. Ezt legjobban egy céltáblához lehet hasonlítani, ahol az alacsony szórással rendelkező modellek egy irányba céloztak és a céltábla egy adott szegletét találták el az összes nyíllal (ha a céltábla közepébe céloztak, akkor mind ahhoz közel landolt, ellenkező esetben egyik sem), vagy magas szórás esetén a nyilak a céltábla minden szegletébe betaláltak, esélyt adva arra, hogy a céltábla közepébe landoljanak. A modellek szórásának értelmezése nem más, mint a „bias-variance tradeoff” paradigma interpretálása a mi esetünkre (NG 2018). A „bias” mutató jelenti a gépi tanuló modell tanulómintán tett teljesítményétől és az elvárt teljesítmény közötti különbséget, míg a modell varianciája a tanulómintán és a tesztmintán értékelt különbségeket fejezi ki a korábban ismertettek alapján. A kettő egyidejű minimális jelenléte vezet modelljószághoz, azonban az említett „bias-variance tradeoff” elmélete alapján egyikről le kell mondani a másik érdekében. A jelen problémában azonosított modellszórás értelmezhető tehát mindezek fényében a gépi tanuló modellek esetén definiált varianciának, azaz elméletileg a cél ennek minimalizálása.

Ugyanakkor, alacsony szórást produkálhatnak a kevésbé ideális modellek is, ezért vélhetően a szórás-irány optimummal rendelkeznek. Mindenképp alapvető az említett felvetés gyakorlati úton történő bizonyítása is, ezért a hasonlóságelemzés közreműködésével, hasonlóan a korábban eljáratokhoz, elvégeztem az irány-preferencia meghatározására irányuló elemzést. A következő táblázat szemlélteti a hasonlóságelemzés által meghatározott értékeket (36. táblázat), ahol a hasonlóságelemzés által becsült fiktív célváltozó értéke 1000 volt, továbbá a modellek érvényességét. A táblázat 2 irányt szemléltet. Érvényesnek tekinthető minden olyan modell, melyre teljesül a függvény-szimmetria, azaz a célváltozó alulról, illetve felülről történő becslése a meghatározott konstans célváltozó középpontjával.

36. táblázat: Felügyelet nélküli modellek leíró statisztikáinak irány-preferencia értékelése hasonlóságelemzéssel

Szórás irány-preferenciák				
ID	Modell azonosító	Minél kisebb, annál jobb	Minél nagyobb, annál jobb	Érvényesség
1	R_NS_NRS_C	997	1008	Érvényes
2	R_NS_NRS_E	1012	999	Érvényes
3	R_NS_NRS_P	1002	1023	Nem érvényes
4	R_NS_RS_C	997	1008	Érvényes
5	R_NS_RS_E	991	1002	Érvényes
6	R_NS_RS_P	1002	1017	Nem érvényes
7	R_S_NRS_C	996	996	Nem érvényes
8	R_S_NRS_E	1003	999	Érvényes
9	R_S_NRS_P	1008	1002	Nem érvényes
10	R_S_RS_C	1002	975	Érvényes
11	R_S_RS_E	991	996	Nem érvényes
12	R_S_RS_P	1002	981	Érvényes
	Érvényes modellek száma (db)	4	3	
	Átlagok	1000	1000	

Forrás: Saját szerkesztés

Alapvetően az állapítható meg a szórás irány-preferenciák hasonlóságelemzés által levezetett kiértékeléséből, hogy feltételezhetően nem szignifikáns a különbség a két szórás irány-preferencia nézet között, melyet a varianciaelemzés is megerősített (20. számú melléklet), mindamellet hajszállal a „minél kisebb a szórás, annál jobb” irány-preferencia a vélelmezett. Itt jelenik meg tökéletesen a bírói ítélkezési logika, azaz az ügyvéd és az ügyész is képes volt alátámasztani az érveit a felmutatott evidenciák által, ezért érdemes mind a két esetet alaposan kivizsgálni, mielőtt elhamarkodottan döntenénk.

Ennél a pontnál kiemelendő, a teljes dolgozatra érvényes végtelen elemzés-potenciál elvének létezése, vagyis az, hogy a mindvégig fekete-fehér erőterként alkalmazott validitás fuzzy-jellegű erőterekké konvertálható, melyekkel újabb műveletek végezhetők, mint az érvényes/nem érvényes állapotokkal (PITLIK L. – PITLIK M. 2021b). A végtelenség ez esetben a modellezési alapvetések tetszőleges finomhangolását jelenti. A dolgozat más részeiben (pl. a „jó” fogalmának operacionalizálásakor, valamint a 4.3 fejezetben, ahol a legjobb modell keresése zajlik, ott a

mindenkori konzisztens zárómodell önmagában is n -féle lehet, ami felveti, hogy az n alternatíva közül melyik a legjobb. Így a „modelleket rangsoroló modelleket rangsoroló modellek rangsorolása” nem a vége a végtelen egymásba ágyazódó modellezési rétegnek. A dolgozat a végtelenre nem fókuszál, de említeni köteles, hogy a végtelen egyelőre emberi döntéssel, illetve a valós idejű elvárások mentén történik minden esetben.

4.4.7. Eredmények értékelése a „minél kisebb, annál jobb” irány-preferencia nézetben

A következő, 37. táblázat szemlélteti az első szórás irány-preferencia („minél kisebb, annál jobb”) nézet szerint az aggregált jóságleíró mátrixból kalkulált modellrangsort, amely meghatározza a verseny győztesét (zölddel jelölve) a jóságleíró mátrix attribútumainak naiv átlagolásával és hasonlóságelemzés által előállított módokon.

37. táblázat: Felügyelet nélküli modellértékelés összefoglaló táblázata az első szórás irány-preferencia nézetben

ID	Modell azonosító	Megállapítások súlyozása	Rendszerek súlyozása	Távolság-metrika	Naiv Átlagok	Naiv rangsor	Hasonlóság-elemzés idealitások	Hasonlóság-elemzés rangsor
1	R_NS_NRS_C	Nincs	Nincs	Koszinusz	5.75	6	997	8
2	R_NS_NRS_E	Nincs	Nincs	Euklideszi	3.00	1	1012	1
3	R_NS_NRS_P	Nincs	Nincs	Pearson	5.75	6	1002	4
4	R_NS_RS_C	Nincs	Súlyozott	Koszinusz	5.75	6	997	8
5	R_NS_RS_E	Nincs	Súlyozott	Euklideszi	4.75	3	991	11
6	R_NS_RS_P	Nincs	Súlyozott	Pearson	5.75	6	1002	4
7	R_S_NRS_C	Súlyozott	Nincs	Koszinusz	6.00	10	996	10
8	R_S_NRS_E	Súlyozott	Nincs	Euklideszi	3.75	2	1003	3
9	R_S_NRS_P	Súlyozott	Nincs	Pearson	5.00	5	1008	2
10	R_S_RS_C	Súlyozott	Súlyozott	Koszinusz	6.25	11	1002	4
11	R_S_RS_E	Súlyozott	Súlyozott	Euklideszi	4.75	3	991	11
12	R_S_RS_P	Súlyozott	Súlyozott	Pearson	6.25	11	1002	4

Forrás: Saját szerkesztés

A táblázat alapján az olvasható le, hogy az R_NS_NRS_E a legideálisabb modell a felvetett gyanúgenerálási probléma megoldására az összes fejlesztett modell között, melyet mind a két rangsor (naiv átlagolás és hasonlóságelemzés is) alátámaszt. Az R_NS_NRS_E modell nem alkalmaz sem megállapítás-súlyozást, sem rendszersúlyozást, távolság/kapcsolat-metrikáját az Euklideszi távolság alapján kalkulálja. A 21. számú melléklet részletezi a mátrix súlyértékeit lépcsőként rendezve.

A következő táblázat részletezi az egyes modell tulajdonságok szerinti csoportosítást és az idealitás mutatók átlagát, mely további alátámasztással szolgálhat az R_NS_NRS_E győzelmének létjogosultságára (38. táblázat).

38. táblázat: Felügyelet nélküli modellek értékeinek összefoglaló táblázata

Átlagok	Megállapítások súlyozása		Rendszerek súlyozása		Távolságmérika		
	Nincs	Súlyozott	Nincs	Súlyozott	Euklideszi	Pearson	Koszinusz
<i>Naiv</i>	5.13	5.33	4.88	5.58	4.06	5.69	5.94
<i>Hasonlóságelemzés</i>	1000.00	1000.00	1002.75	997.25	999.25	1003.00	997.75

Forrás: Saját szerkesztés

A táblázat alapján eldönthető, hogy a naiv átlagok alapján létjogosultságunk van az R_NS_NRS_E modellt „megkoronázni”, mivel az összes kategória szerint az az ideálisabb modell, azonban a hasonlóságelemzés szerint a megállapítások súlyozása között nincs eltérés, az alkalmazott távolság/kapcsolat-metrika tekintetében az euklideszi modellek csak 2. helyezést értek el.

A varianciaelemzés arra mutat rá, hogy a hasonlóságelemzés F-próba értékei nem szignifikánsak, a naiv átlagolás esetében csak a távolságmérika F-próba értéke mutat szignifikáns értéket (szóráshomogenitás feltételének teljesülése nélkül), mely eredményeket a 22. számú melléklet összegez. Ez azt kívánja vélelmezni, hogy a hasonlóságelemzés által, bár van különbség az egyes modellek között (mivel értékeik 1000-tól különböznek), azok nem szignifikánsak. A modelljóságok értelmezése így, további kérdéseket vethet fel, mely jelenség minél kisebb az adatvagyon, annál triviálisabb, hogy fel kell, hogy lépjen a megoldások alternatívitásainak elkerülhetetlen végtelensége okán.

Elvégezve a teszthalmoz kiértékelését a szimulált környezetben és összehasonlítva az eredményeket a leírtakkal, a következő táblázat szemlélteti a modellek találati arányaira vonatkozó összefoglaló eredményeket (39. táblázat).

39. táblázat: Felügyelet nélküli modellek értékelő táblázata

ID	Modell azonosító	Igaz pozitívak száma (db)	Igaz negatívak száma (db)	Hamis pozitívak száma (db)	Hamis negatívak száma (db)	Pontosság	Precizitás	Fedés	F1-Pont
1	R_NS_NRS_C	5	183	45	7	0.78	0.12	0.42	0.14
2	R_NS_NRS_E	8	193	35	4	0.84	0.24	0.67	0.32
3	R_NS_NRS_P	5	195	33	7	0.83	0.13	0.42	0.16
4	R_NS_RS_C	5	183	45	7	0.78	0.12	0.42	0.14
5	R_NS_RS_E	11	136	92	1	0.61	0.22	0.92	0.30
6	R_NS_RS_P	5	195	33	7	0.83	0.13	0.42	0.16
7	R_S_NRS_C	5	175	53	7	0.75	0.11	0.42	0.13
8	R_S_NRS_E	8	186	42	4	0.81	0.22	0.67	0.30
9	R_S_NRS_P	5	186	42	7	0.80	0.12	0.42	0.14
10	R_S_RS_C	5	172	56	7	0.74	0.11	0.42	0.13
11	R_S_RS_E	11	123	105	1	0.56	0.18	0.92	0.25
12	R_S_RS_P	5	184	44	7	0.79	0.12	0.42	0.15

Forrás: Saját szerkesztés

Az euklideszi modellek találták el legnagyobb számban helyesen a gyanúsnak megjelölt kontrollokat, azaz a legmagasabb az igaz pozitív találatok száma, mely a Precizitás, Fedés és F1 pontszám mutatókban is tetten érhető. Továbbá, az R_NS_NR_E modell volt az, amely az igaz negatívakat a többi modellhez képest is relatív kiemelkedően megtalálta (egyedül az R_NS_NRS_P és R_NS_RS_P modellek előzték meg 2 db találattal), ezért valóban járhat neki a kitüntetett első helyezés. Megtekintve a rendszersúlyozott euklideszi modelleket, az látható, hogy igen alacsony igaz negatív találatokkal bírnak, melyet az objektum leíró tulajdonságoknál többek között a rendszerválaszok száma (T02) volt képes rögzíteni és ez által büntetni a modelleket, ennek okán érthetővé válik, hogy miért nem az euklideszi modellek lettek „csapatversenyben a győztesek”, a legrosszabb igaz negatív találati aránnyal rendelkeztek és ezt a súlyozások sem voltak képesek kompenzálni. Ez azt jelenti, hogy míg bizonyos konfigurációval nyertek az euklideszi távolságmetrikát alkalmazó modellek, azt elveszítették egy másikban, amit az alacsony modell-idealitás-értékek is alátámasztanak. A leírtakat összegezve, az következik, hogy a hasonlóságelemzés által produkált eredmények magas magabiztossággal alátámasztják a függetlennek vélt performancia metrikákat.

Felmerül, azonban az a kérdés, hogy miért nem tudott egyöntetűen dönteni a hasonlóságelemzés abban az esetben, amikor arra a kérdésre kerestük a választ, hogy a súlyozott megállapítású vagy a nem súlyozott megállapítású modellek voltak-e a jobbak? Az alábbi táblázat szemlélteti a performancia metrikák értékeit kategóriánkénti megbontásban, amely választ ad a kérdésre (40. táblázat).

40. táblázat: Felügyelet nélküli modellek összefoglaló táblázata

Jóságmutatók	Megállapítássúlyozás		Rendszersúlyozás		Távolság/kapcsolat-metrika		
	Nincs	Súlyozott	Nincs	Súlyozott	Euklideszi	Koszinusz	Pearson
<i>Igaz pozitívak száma (db)</i>	6.50	6.50	6.00	7.00	9.50	5.00	5.00
<i>Igaz negatívak száma (db)</i>	180.83	171.00	186.33	165.50	159.50	178.25	190.00
<i>Hamis pozitívak száma (db)</i>	47.17	57.00	41.67	62.50	68.50	49.75	38.00
<i>Hamis negatívak száma (db)</i>	5.50	5.50	6.00	5.00	2.50	7.00	7.00
<i>Pontosság</i>	0.78	0.74	0.80	0.72	0.70	0.76	0.81
<i>Precizitás (Precision)</i>	0.16	0.14	0.15	0.14	0.21	0.11	0.12
<i>Fedés (Recall)</i>	0.54	0.54	0.50	0.58	0.79	0.42	0.42
<i>F1 pontszám</i>	0.20	0.18	0.20	0.19	0.29	0.14	0.15

Forrás: Saját szerkesztés

A hasonlóságelemzés nem talált különbséget a megállapítássúlyozás alkategóriái között, mivel számos mutató értéke megegyezik (pl. Igaz pozitívak száma, Hamis negatívak száma, Fedés stb.), vagy sok esetben csak minimálisan különbözik a két érték egymástól (pl. Pontosság, Precizitás, F1 pontszám stb.), ezért a hasonlóságelemzés az eltéréseket nem értékelte úm. „szignifikánsnak” az objektumleíró tulajdonságok alapján.

4.4.8. Eredmények értékelése a „minél nagyobb, annál jobb” irány-preferencia nézetben

Amennyiben a magasabb szórású modelleket önkényesen jobbnak akarjuk tekinteni, a 41. táblázat alapján láthatóvá válik, hogy abban az esetben az ideális modell a gyanúgenerálásra az R_NS_NRS_P lesz.

41. táblázat: Felügyelet nélküli modellértékelés összefoglaló táblázata a második szórás irány-preferencia nézetben

ID	Modell azonosító	Megállapítások súlyozása	Rendszerek súlyozása	Távolság-metrika	Naiv Átlagok	Naiv rangsor	Hasonlóság-elemzés idealitások	Hasonlóság-elemzés rangsor
1	R_NS_NRS_C	Nincs	Nincs	Koszinusz	4.25	3	1008	3
2	R_NS_NRS_E	Nincs	Nincs	Euklideszi	5.75	7	999	7
3	R_NS_NRS_P	Nincs	Nincs	Pearson	3.00	1	1023	1
4	R_NS_RS_C	Nincs	Súlyozott	Koszinusz	4.25	3	1008	3
5	R_NS_RS_E	Nincs	Súlyozott	Euklideszi	6.00	8	1002	5
6	R_NS_RS_P	Nincs	Súlyozott	Pearson	3.50	2	1017	2
7	R_S_NRS_C	Súlyozott	Nincs	Koszinusz	5.25	6	996	9
8	R_S_NRS_E	Súlyozott	Nincs	Euklideszi	6.00	8	999	7
9	R_S_NRS_P	Súlyozott	Nincs	Pearson	4.75	5	1002	5
10	R_S_RS_C	Súlyozott	Súlyozott	Koszinusz	7.00	12	975	12
11	R_S_RS_E	Súlyozott	Súlyozott	Euklideszi	6.50	10	996	9
12	R_S_RS_P	Súlyozott	Súlyozott	Pearson	6.50	10	981	11

Forrás: Saját szerkesztés

A magasabb szórás azt jelenti, hogy a tesztalmazon végzett predikciók esetén voltak előkelő helyen végzett objektumok (ahol a gyanú megítélése szinte hibátlan volt) és hátrébb végzett objektumok is (ahol a gyanú helyességének megítélése gyengének nevezhető). A magas szórású modellek kisebb fókusszal rendelkeznek, azonban nagyobb esélyt adnak arra, hogy akár véletlen is, de kifogás nélkül eltalálják a hibásan auditált kontrollokat. A 23. számú melléklet részletezi a mátrix súlyértékeit lépcsőként rendezve. A következő táblázat részletezi az egyes modell tulajdonságok szerinti csoportosítást és az idealitás mutatók átlagát (42. táblázat).

42. táblázat: Felügyelet nélküli modellek értékeinek összefoglaló táblázata

	Megállapítások súlyozása		Rendszerek súlyozása		Távolságmetrika		
Átlagok	Nincs	Súlyozott	Nincs	Súlyozott	Euklideszi	Pearson	Koszinusz
Naiv	4.46	6.00	4.83	5.63	6.06	4.44	5.19
Hasonlóságelemzés	1009.00	991.00	1005.00	997.00	998.50	1005.25	996.25

Forrás: Saját szerkesztés

A naiv átlagolás és hasonlóságelemzés is döntésképes volt ezúttal, azonban a távolságmetrika-alapú kategória esetén a hasonlóságelemzés az Euklideszi-t, míg a naiv átlagolás a Koszinusz

hasonlóság alapú metrikát sorolta a 2. helyre. Az R_NS_NRS_P modell győzelme teljesen beigazolódtott, minden kategóriát megnyert, így nem maradt a „bíróági tárgyaláson” ellenérv ellene.

Mindkét szórás-nézetből az is szembetűnik, hogy a súlyozással ellátott modellek alapvetően kevésbé ideálisabbnak bizonyultak, mely feltételezés is úgy szint beigazolódtott. Az átlagokat összehasonlítva varianciaelemzés (24. számú melléklet) által a megállapítások súlyozásának F-próba értéke a naiv átlagok és hasonlóságelemzés alapján is szignifikánsnak tekinthetők, tehát a preferált szórás-nézetben szignifikánsan jobbak a nem súlyozott megállapítással kalkuláló modellek, melyeket a korábban elemzett 38. összefoglaló táblázatban ismertettek esetén nem állt fent. Amennyiben, tehát a kényszerűen a magasabb szórású modelleket preferáljuk, akkor az R_NS_NRS_P modell a javasolt modell, mely a legjobb igaz negatív találati aránnyal bír, és ezáltal a Pontosság metrikája 0.01-el tér el az előző szórás-nézetben kikiáltott győztesétől.

A Koszinusz hasonlósággal ellátott modellek egyik szórásnézetben sem tudtak győzni, mely feltehetően köszönhető annak, hogy egyik nézetben sem volt kiugróan jó értékeik, mely a szűkített adathalmaz sajátosságainak tudható be. Míg az Euklideszi modellek igaz pozitív találati arányai a legjobbak voltak, a Pearson modellek igaz negatív találati arányaik voltak a kedvezőbbek, tehát a Koszinusz hasonlósággal kalkuláló modellek a kettő modell eredményei között helyezkedtek el. A 24. számú mellékletben látható táblázatok alapján észrevehető, hogy ezen modellek is magas szórással rendelkeztek, azonban a Pearson modellek rangsorai átlagosan jobbak voltak, amely azt jelenti, hogy az objektumleíró tulajdonságaik közelebb kerültek a kívánatos állapottól.

Mivel a hasonlóságelemzés sem tudott túlnyomó többségen dönteni arról, melyik szórás iránypreferencia a jobb, ez így a döntéshozó elvárásainak és erőforrásainak mértékétől függ ismételt, melyet a döntéshozó kockázati étvágya befolyásol.

4.4.9. Objektív modelljóság-becslés a generált modelleken anti-diszkriminatív eljárással

Végezzük el az alkalmazott performancia metrikák (Pontosság, Precizitás, Fedés, F1-Pont) kiértékelését a hasonlóságelemzés által nyújtott anti-diszkriminatív matematikai apparátus által, ahol a norma értéke 1000 (43. táblázat). Kijelenthető, hogy az R_NS_NRS_E modell mutatói összességében a legideálisabbak (1016.60), tehát a visszacsatolás eredménye az, hogy jogosan lett a modell győztesként kihirdetve. Az R_NS_NRS_P modell csak a harmadik helyezést ért el (az R_NS_RS_P modellel egyidejűleg), mely azt vélelmezi, hogy az első szórás iránypreferencia („minél kisebb, annál jobb”) volt az ideális irány.

43. táblázat: Modelljóság becslés anti-diszkriminatív eljárással

ID	Modell azonosító	Pontosság	Precizitás	Fedés	F1-Pont	Y ₀
1	R_NS_NRS_C	0.78	0.12	0.42	0.14	995.60
2	R_NS_NRS_E	0.84	0.24	0.67	0.32	1016.60
3	R_NS_NRS_P	0.83	0.13	0.42	0.16	1005.60
4	R_NS_RS_C	0.78	0.12	0.42	0.14	995.60
5	R_NS_RS_E	0.61	0.22	0.92	0.30	1003.60
6	R_NS_RS_P	0.83	0.13	0.42	0.16	1005.60
7	R_S_NRS_C	0.75	0.11	0.42	0.13	986.60
8	R_S_NRS_E	0.81	0.22	0.67	0.30	1011.60
9	R_S_NRS_P	0.80	0.12	0.42	0.14	997.60
10	R_S_RS_C	0.74	0.11	0.42	0.13	985.60
11	R_S_RS_E	0.56	0.18	0.92	0.25	998.60
12	R_S_RS_P	0.79	0.12	0.42	0.15	997.60

Forrás: Saját szerkesztés

A 37. táblázat mutatta be a modellek rangsorolását, mely az objektum-leíró tulajdonságok felhasználásával került levezetésre. Korrelációt számítva a 37. és 43. táblázatban ismertetett hasonlóságelemzés idealitásokkal (Y₀) a korrelációs együttható értéke: 0.42, mely közepes pozitív kapcsolatot feltételez. Kiemelendő, a modell-preferencia mind a két eljárás esetén az R_NS_NRS_E volt, független visszaigazolást nyert a levezetés eredményessége.

Varianciaelemzéssel elvégezve a kategóriák közötti különbségek szignifikanciájának mérését, a 25. számú mellékletben közölt táblázatok alapján megállapítható, hogy az EUC modellek a COS modelleknél szignifikánsan jobbak, míg az EUC és PEA, valamint a PEA és COS modellek között nincs szignifikáns különbség. Az egyes kategóriák között (rendszerűlyozás és megállapítássúlyozás), hasonlóan nincs szignifikáns különbség.

4.4.10. Az alfejezet összefoglalása

Az objektumleíró tulajdonságok és a performancia metrikák közötti összefüggés beigazolódott a 4.4. alfejezetben tárgyalt levezetésekben, tehát valóban a modellektől elvárt, racionálisan létrehozott attribútumok képesek megbecsülni a modelljóságok aggregált értékeit. Megállapítható, hogy:

- A modell-preferencia keresésére fejlesztett eljárás alkalmas az ideális modell megtalálására egy előre definiált (véletlenszerű, célzott) modellhalmazból, mely független teszteléssel objektíven bizonyítható;
- A modellek önerősítő mechanizmusai (rendszerválaszok viselkedései) és modelljóság metrikák között összefüggés tapasztalható, így az ellentmondásosság felfedezése révén a modellek teljesítményei becsülhetővé válnak;
- A döntéstámogató rendszer által közölt rendszerválaszokhoz tartozó attribútumok iránypreferenciának meghatározása objektivizálható, így lehetséges a döntéselőkészítési folyamatban minimalizálni az emberi tévesztéseket, valamint tudatosítani lehet az emberi önkényességet.

A leírtak alapján kijelenthető, hogy az alfejezet elején ismertetett hipotézis igazolható – új megoldást jött létre a korábban létező megoldások, mint benchmark-ok halmazához.

A problémák végtelen egymásra rétegződésének speciális nézete, vagyis az éles tesztelés sikerességének termelési függvénye nem része jelen dolgozatnak, de ennek létre kényszerűen utalni kell annak érdekében, hogy a komplexitás (kutatási feladat) következő szintje (jövőképe is) értelmezhető legyen a dolgozattal szemben támasztott elvárásoknak is megfelelően.

4.5. A dolgozat célkitűzéseinek teljesítése a SMART feltételrendszer alapján

A kutatás célkitűzéseinek értékelését a dolgozat 1. fejezete ismertette, melyet az 1. táblázat összegzett a SMART feltételrendszer alapján. A 44. táblázat szemlélteti a kutatás célkitűzéseinek teljesítését az 1. táblázattal összhangban.

44. táblázat: A kutatás célkitűzéseinek teljesítése a SMART feltételrendszer alapján

Kritériumok	C1	C2	C3
<i>Tényleges</i>	Információbiztonsági kontrollhiányosságok detektálására irányuló döntéstámogató rendszer fejlesztése (robot-auditor)	Genetikai potenciál-alapú új keresésvezérlés iterációként változtatható irány-preferenciákkal (új funkcionalitás)	Klasszikus tesztelés nélküli modell-preferencia levezetés objektív és automatizált aggregált jószág fogalom alapján (új funkcionalitás)
<i>Mérhető</i>	Objektív jószágmetrikák alkalmazása, melyek kifejezik a döntéstámogató rendszer idealitását (pl. F1-Pont, AUROC stb.)	Objektív jószágmetrika alkalmazása, mely kifejezi a kereső eljárás ideális irányba történő elmozdító hatását (F1-Pont)	Objektív jószágmetrikák alkalmazása, melyek kifejezik a modell-preferencia levezetés idealitását (pl. hasonlóságelemzés idealitás mutató, F1-Pont stb.)
<i>Teljesíthető</i>	A 4.2. alfejezet bizonyította a megfogalmazott célkitűzés teljesítését	A 4.3. alfejezet bizonyította a megfogalmazott célkitűzés teljesítését	A 4.4. alfejezet bizonyította a megfogalmazott célkitűzés teljesítését
<i>Releváns</i>	A 2.2. alfejezetben ismertetett szakirodalmi áttekintés bizonyította a célkitűzés relevanciáját	A 2.3. alfejezetben ismertetett szakirodalmi áttekintés bizonyította a célkitűzés relevanciáját	A 2.3. alfejezetben ismertetett szakirodalmi áttekintés bizonyította a célkitűzés relevanciáját
<i>Időhöz kötött</i>	Az előre definiált határidőre történő leszállítás megtörtént	Az előre definiált határidőre történő leszállítás megtörtént	Az előre definiált határidőre történő leszállítás megtörtént

Forrás: Saját szerkesztés

5. ÚJ ÉS ÚJSZERŰ TUDOMÁNYOS EREDMÉNYEK

Kutatásom eredményeül az alább felsorolt új és/vagy újszerű tudományos eredményeket rögzítettem a hipotézisekkel összhangban:

H1: *Az információbiztonsági auditjelentések szöveges eredményeiből strukturált adatbázist alkotva és bemenetként a mesterséges intelligencia fogalmkörébe illeszthető eszközökkel azt feldolgozva, az auditok során feltárni kívánt kontrollhiányosságok megléte a véletlen találgatásnál nagyobb valószínűséggel kimutathatók, azaz a kontrollhiányosságok konstellációi matematikailag értelmezhető összefüggéseket hordoznak magukban.* **IGAZOLT**

Bizonyítás típusa: új, saját fejlesztésű innovatív döntéstámogató rendszer (robot-auditor) fejlesztése információbiztonsági kontrollhiányosságok detektálására.

- Az audit által közölt megállapítások, azaz kontrollhiányosságok, melyek szöveges riport formájában kerülnek közlésre a szervezetek vezetői és befektetői számára, logikus kontextusfüggő formában kategorizálhatók, a megállapítások tematikájára irányuló tudatos rendszerezést követően a gép számára strukturált formában átadható (4.2.1. alfejezet);
- Az auditok által dokumentált megállapítások együttes konstellációi között összefüggés tapasztalható mely objektíven, matematikai apparátusok felhasználásával kimutatható, így döntéstámogató rendszer építhető, mely képes az auditok hatékonyságát növelni azok objektív minőségbiztosítása révén (4.2.6. alfejezet):
 - Egyszerű és hibrid ABM F1-Pont értékei rendre: 0.57 és 0.67;
 - Egyszerű és hibrid GBM F1-Pont értékei rendre: 0.43 és 0.55;
 - Egyszerű és hibrid NN F1-Pont értékei rendre: 0.53 és 0.56;
- A kontrollhiányosságok együttes megléte alapján előre a döntéstámogató rendszer által nem feldolgozott adatokon a kontrollhiányosságok adott kontrollra becsülhetők, a véletlen találgatásnál ($AUROC > 0.5$ és $AUPRC > 0.14$) ideálisabb eredmény érhető el (4.2.6. alfejezet):
 - Egyszerű és hibrid ABM AUROC értékei rendre: 0.85 és 0.90, AUPRC értékei rendre: 0.58 és 0.70;
 - Egyszerű és hibrid GBM AUROC értékei rendre: 0.85 és 0.85, AUPRC értékei rendre: 0.61 és 0.71;
 - Egyszerű és hibrid NN AUROC értékei rendre: 0.82 és 0.85, AUPRC értékei rendre: 0.56 és 0.62;

H1.1: *A gyanúgenerálás, mint megoldandó üzletileg értelmezett probléma sajátosságait értékelve, a kontrollhiányosságok detektálása megoldható felügyelt és felügyelet nélküli gépi tanuló eljárásokkal is.* **IGAZOLT**

- A kontrollhiányosságok detektálását osztályozási problémaként azonosítva, a gyanúgenerálás lehetséges felügyelt gépi tanuló modellezés keretében megvalósítani, ahol a dedikált célváltozó az adott kontroll megfelelőségére irányuló állapotot jelöli (4.2.4. alfejezet):
 - Egyszerű ABM F1-Pont értéke: 0.57;
 - Egyszerű GBM F1-Pont értéke: 0.43;
 - Egyszerű NN F1-Pont értéke: 0.53;

- A kontrollhiányosságok detektálását ajánlórendszerként azonosítva, a gyanúgenerálás lehetséges felügyelet nélküli gépi tanuló modellezés keretében megvalósítani, amely célváltozó hiányában, egy listával tér vissza az azonosított gyanúmomentumokról, melyet az auditjelentések hasonlósága alapján vél felfedezni (4.2.5. alfejezet):
 - EUC F1-Pont értéke: 0.21;
 - PEA F1-Pont értéke: 0.21;
 - COS F1-Pont értéke: 0.22;

H1.2: *A gyanúgenerálás teljesítménye fokozható hibrid megközelítésben, azaz a felügyelt és nem felügyelt módszerek együttes felhasználásának a kutatásban alkalmazott releváns performancia metrikái ideálisabb értékeket mutatnak, mint önálló alkalmazásban. IGAZOLT*

- A kontrollhiányosságok detektálásának teljesítménye a felügyelt és felügyelet nélküli módszerek együttes alkalmazásával javítható, ahol a felügyelet nélküli ajánlórendszer javaslatot tesz a gyanúmomentumokról, melyet a felügyelt módszerek beépítenek a döntéselőkészítésbe, így addicionális információként csökken a bizonytalanság mértéke (4.2.6. alfejezet):
 - Egyszerű és hibrid ABM F1-Pont értékei rendre: 0.57 és 0.67;
 - Egyszerű és hibrid GBM F1-Pont értékei rendre: 0.43 és 0.55;
 - Egyszerű és hibrid NN F1-Pont értékei rendre: 0.53 és 0.56;

H1.3: *A hibrid modell többlet-információs értéket teremtve képes az egyszerű modellek általánosító képességén javítani. RÉSZBEN IGAZOLT*

- A felügyelt gépi tanuló rendszerek varianciáit csökkenti a hibrid megközelítés, ezért általánosító képességük ideálisabb, mint önálló alkalmazásban (4.2.6 alfejezet):
 - Egyszerű és hibrid ABM Variancia értékei rendre: 0.08 és 0.06;
 - Egyszerű és hibrid GBM Variancia értékei rendre: 0.04 és 0.03;
 - Egyszerű és hibrid NN Variancia értékei rendre: 0.11 és 0.11 (a neurális háló esetén nem csökkent a variancia);

H2: *A döntéstámogató rendszer genetikai potenciálja letapogatható hasonlóságelemzéssel ellátott kereső eljárással a tanításra alkalmazott adathalmaz irányított feldolgozásán keresztül, úgy, hogy a genetikai potenciálhoz vezető kereső eljárás a genetikus algoritmusok esetén alkalmazott véletlen mutáció és a populáció egyedeinek keresztezése nélkül is képes ideálisabb eredményt szolgáltatni. IGAZOLT*

Bizonyítás típusa: új, saját fejlesztésű innovatív keresési eljárás.

- A hasonlóságelemzés által generált lépcsős függvények alkalmasak keresési eljárásokban történő felhasználásra (4.3.2. alfejezet);
- A döntéstámogató rendszer genetikai potenciálja kereshető a tanulóhalmazhoz tartozó racionális leíró attribútumok érték-irány levezetésével (4.3.2. alfejezet);

- A keresési eljárás meg tudja határozni a rendelkezésre álló populáció értékei alapján, hogy adott modell vélhetően elérte-e már a tanulási halmaz optimalizálási kísérlete révén genetikai potenciálját, valamint, hogy milyen attribútumok módosítása szükséges az ideális célváltozó irányába történő elmozduláshoz (4.3.2. alfejezet):
 - Az algoritmus a legelőkelőbb helyen álló súlyszámok összegeként megadja a modell genetikai potenciálját adott célváltozó tükrében (a GBM esetében a 11. iterációban az F1-Pont genetikai potenciálja: 0.71);
 - Az algoritmus meghatározta a súlyszámok különbségéből adódóan, mely attribútumok módosítása volt szükséges az ideális célváltozó irányába történő elmozduláshoz (pl. a GBM esetében az 1. iterációban az f_{10} attribútum meghatározott irány-preferencia szerinti növelésével lehetett elérni, ahol annak értéke 1.9 volt).
- A kereső eljárás véletlen mutáció nélkül is képes ideálisabb irányt megjelölni (4.3.2. alfejezet):
 - Az algoritmus véletlen mutáció nélkül már az 1. iterációban javította a GBM F1-Pont értékét 0.60833-ról 0.61925-ra, mely 1.80 %-os javulást jelentett;
 - Az algoritmus véletlen mutáció nélkül a 11. iterációban a GBM F1-Pont értékét a kezdeti 0.60833-ról 0.63673-ra javította, mely 4.57 %-os javulást jelentett.
- A kereső eljárás az egyedek keresztezése nélkül is képes ideálisabb irányt megjelölni (4.3.2. alfejezet):
 - Az algoritmus az egyedek keresztezése nélkül már az 1. iterációban javította a GBM F1-Pont értékét 0.60833-ról 0.61925-ra, mely 1.80 %-os javulást jelentett.
 - Az algoritmus az egyedek keresztezése nélkül a 11. iterációban a GBM F1-Pont értékét a kezdeti 0.60833-ról 0.63673-ra javította, mely 4.57 %-os javulást jelentett.
- A kereső eljárás nem igényel a nyitó populáción túl köztes/iterációnként új populációkat (4.3.2. alfejezet);
- Az irányok automatikusan minden iterációban újra paraméterezhetők (4.3.2. alfejezet).

H3: *A mesterséges intelligenciával ellátott döntéstámogató rendszerek teljesítményalapon a gépi tanuló alkalmazások klasszikus tesztelési eljárásai nélkül is rangsorolhatók, a predikciók, mint generált gyanúforrások leíró tulajdonságainak érték-irány levezetésével és az ezen adatokat feldolgozó matematikai apparátussal, mely automatizáltan képes a preferált modellek objektív meghatározására. IGAZOLT*

Bizonyítás típusa: új, saját fejlesztésű innovatív eljárás.

- A modell-preferencia keresésére fejlesztett eljárás alkalmas az ideális modell megtalálására egy előre definiált modellhalmazból, mely független teszteléssel objektíven bizonyítható (4.4.7. és 4.4.9. alfejezetek):
 - Az R_NS_NRS_E modell volt a legideálisabb, melynek hasonlóságelemzés idealitása: 1012 (a legmagasabb az összes modell között), melyet alátámasztott a függetlenül mért modelljóság becslés (a hasonlóságelemzés idealitása: 1017, a legmagasabb az összes modell között).

- A modellek önerősítő (konzisztencia) mechanizmusai (rendszerválaszok viselkedései) és modelljóság metrikák között összefüggés tapasztalható, így az ellentmondásosság felfedezése révén a teljesítmény becsülhetővé válik (4.4.9. alfejezet):
 - Korrelációt számítva a hasonlóságelemzés idealitásokkal (Y_0) a korrelációs együttható értéke: 0.42, mely közepes pozitív kapcsolatot feltételez.
- A döntéstámogató rendszer által közölt rendszerválaszokhoz tartozó attribútumok iránypreferenciának meghatározása objektivizálható, így lehetséges a döntéselőkészítési folyamatban minimalizálni az emberi tévesztéseket (4.4.5 alfejezet).

6. KÖVETKEZTETÉSEK ÉS JAVASLATOK

A dolgozatban leírtak alapján az alábbi pontok összegzik a kutatás és kísérletek által nyert gyakorlati alkalmazhatóság területeit:

- Audit, és auditot végző szakmai területek a dolgozatban közölt eljárásokat alkalmazni tudják, mely többek között hozzájárul a biztonságosabb informatikai üzemeltetéshez, jogszabályozói elvárások betartásához, valamint növelheti a szervezet befektetőinek bizalmát;
- A kontrollhiányosságok automatizált detektálása képes pl. az emberi hibából fakadó tévesztések felülvizsgálatára és felülbírálatára, objektívebb képet ad az ellenőrzés minőségéről befolyásmentesen (nem fél a következményektől);
- Az információbiztonsági környezet javítása révén a szervezetek kisebb kitettséget mutathatnak külső támadókkal, továbbá a munkatársak szándékos károkozásával szemben, mely növelheti a vevők a szervezet termékei/szolgáltatásai iránti bizalmát, tehát csökkenti a reputációs és kártérítési kockázatokat;
- Az audit megállapítások pontosításának révén leleplezhetővé válhatnak a csalások, szűk keresztmetszetek, mérsélkelhető a vállalati adatvagyon bizalmasságának, integritásának és rendelkezésre állásának kompromittálódásából fakadó kockázatok;
- Az ismertett gépi tanuló rendszerek genetikai potenciálját kereső algoritmus általánosítható, nem kizárólag kontextusfüggő módon működik, hanem egyéb üzleti és szakmai problémák gépi tanúlással feldolgozott megoldáskeresésére is felhasználható;
- A kereső eljárás a tanulóhalmaz méretének/összetételének optimalizálását célozta meg, azonban belátható módon, adott célváltozó esetén képes a célváltozóhoz köthető független változók elmozdulásai alapján, illetve azok érték-irány levezetésével ideálisabb irány kijelölésére, így jobb célváltozó-értékek lépésről lépésre történő megtalálására, valós/sztokasztikusan működő input-output-rendszerek esetében;
- A definiált populációból történő modell-preferencia matematikai levezetése és az ideális modell keresése hasonlóan általánosítható, így lehetőség nyílik kontextus független módon az algoritmust felhasználni ott, ahol eddig a hermeneutikai alrendszer az emberi intuícióra hagyatkozott a több értékelési réteg aggregálása érdekében – így az automatizmus egyszerre garantálja az optimalizálást és a Knuth-i elv betartását;
- Az irány-preferenciák automatizálása által csökkenthető az emberi tévesztés kockázata.

A kutatás során számos olyan feladat és potenciális kutatási irány fogalmazódott meg bennem, melyet szerettem volna elvégezni és dokumentálni, azonban idő és tartalmi korlátok révén nem volt rá lehetőség. Az alábbi pontok egyfajta jövőkép jelleggel foglalják össze azon javasolt kutatási irányokat, mellyel a dolgozatban közölt eredmények javíthatók, valamint a dolgozat alapot szolgált további kutatási célkitűzések meghatározásához:

- A dolgozat a kontrollhiányosságok detektálását osztályozási problémaként azonosítottam, azonban regressziós algoritmusokkal lehetőség adódhat a kontrollokhoz tartozó megállapítások számának approximációjára is;
- A kontrollhiányosságok becslésére az üzleti probléma bizalmasságának sajátosságai miatt, nem állt rendelkezésre számos olyan attribútum, mely képes lenne pontosítani az előrejelzéseket. Amennyiben van rá lehetőség, javasolt a kísérlet megismétlése, ahol elérhető az egyes szervezetek árbevétel, információbiztonsági költségvetés, szervezeti hierarchia, méret stb. adatai, vélelmezhetően javítana a rendszerek ítélőképességén;
- A kutatásban nem került elkülönítésre az egyes auditesztekhez köthető tervezet, implementáció és működési hatékonysági szinten mért kontrollhiányosságok, ezért a

döntéstámogató rendszer képes lehet ezen információk finomhangolásával ítélni, vajon kontrollhiányosság tapasztalható már a kontroll kialakítási, implementációs vagy működtetési fázisaiban;

- A kutatásban strukturált, numerikus értékekkel feltöltött adatbázis került feldolgozásra, azonban a szövegesen megfogalmazott riportok további információt is tartalmazhatnak, mely segítheti az előrejelzések pontosságát, ezért javasolt egy olyan kísérletet is elvégezni, amely ötvözi a természetes nyelvi feldolgozás előnyeit a dolgozatban ismertetett módszerekkel;
- A felügyelet nélküli ajánlórendszer a rendelkezésre álló minta (127) nagyságát szem előtt tartva nem használt fel komplex struktúrákat, melyek a dolgozat írásának időpontjában javasol a szakirodalom pl. mátrix faktorizáció, mely javíthat az ajánlások pontosságán és „perszonalizáltságán”. Feltehetően a mátrix faktorizációs eljárás nem hozott volna többlet javulást jelen adatvagyon számára, mivel a módszer kiaknázása több adatot igényel;
- A kutatás nem kísérlete meg a tökéletes konfiguráció megtalálását, nem volt célja tartalmi korlátok miatt. Javasolt a konfigurációk optimális beállítását követően kiértékelni az algoritmusokat. Továbbá, a konfigurációk előkelőbb megtalálásához a dolgozatban ismertetett kereső eljárás alkalmasnak bizonyul;
- Javasolt a genetikai potenciál kereső eljárás továbbfejlesztése egyidejűleg több célváltozó feldolgozására;
- A hasonlóságelemzés teljesítményét növelendő, javasolt az algoritmus fejlesztése univerzális approximátorok pl. neurális háló által. Az anti-diszkriminatív modellek efféle javítása alkalmassá teheti őket a normától történő eltérés pontosításában;
- Javasolt az irány-preferenciák objektív meghatározása érdekében további algoritmusokkal kísérletezni, mely ideálisabb eredménnyel szolgálhat az irány-preferenciákat illetően, akár hibrid megoldásban képes a „hazudós” modellek és függvények leleplezésére.

A kutatásban ismertetett eredmények és a döntéstámogató rendszer (robot-auditor), véleményem szerint, egy olyan gyakorlati eszközt képes adni az auditorok, IT üzemeltető és biztonsági vezetők, belső jogászok, stb. kezébe, mellyel objektív, független megfelelést és IT biztonsági kontrollkörnyezet kiépítését lehet támogatni és fenntartani, mely valós értéket képez a piaci szereplők részére. Valós piaci kereslet mutatkozik az automatizált biztonságot fokozó megoldások (pl. adatszivárgás-megelőző rendszerek) iránt, melyet a jogszabályi követelmények pl. GDPR még inkább kikényszerít. Mindamellett, a dolgozatot nem is lett volna érdemes elkészíteni, ha nem tartalmazna az ismertetett robot olyan elemeket és funkciókat, melyekre nem mutatkozna piaci kereslet.

Mivel a belső ellenőr (auditor) egy adott vállalat belső alkalmazottja, ezért a vállalat ügyfelei számára kockázatot jelenthet a függetlenségi kérdés, melyet egy objektíven működő, külön műszakilag szeparált (pl. dedikált szerveren) és ellenőrzött (auditált) rendszer mérésélni tud, így potenciális piaci előnyekre, valamint a megnevezett jogszabályi követelmények kényszerítő erejére építkezve, egy startup cég formájában érdemes lehet a megoldást piacosítani. Potenciális célközönség, ezen felül, az auditor cégek, melyek a robot-auditor minőségbiztosítási funkciója révén képesek a reputációs kockázatok csökkentésére.

Gyakorlati megvalósításban, a kutatásban közölt eredmények és maga a forráskód megfeleltethető egy MVP-nek (Minimum Valuable Product - Prototípus), melyhez grafikus interfészt fejlesztve lehetőség adódhat kockázati tőkebefektetők felé történő prezentálásnak.

Műszakilag felhőszolgáltatóra épített infrastruktúrát lenne érdemes kiépíteni az imertetett megoldás mögé (pl. AWS, Azure, IBM Cloud, stb.), mely a skálázhatóság és optimális erőforrás kihasználás miatt a legjobb megoldás egy olyan cég számára, mely akár az első hónapokban robbanásszerűen növekedhet pl. a befektetés volumene, vagy a globális jelenlét miatt elérhető ügyfelek számának növekedése végett.

7. ÖSSZEFOGLALÁS

Az információtechnológia fejlődésének köszönhetően eddig nem látott mértékben történik az üzleti folyamatok automatizálása, mely az ismétlődő manuális feladatok korszerűsítésén túl az emberi hibákban és tévedésekben gazdag munkafolyamatokat szándékozik objektivizálni és hatékonyabbá tenni. Az információtechnológiára irányuló fejlesztési projektek bővelkednek újfajta kockázatokban, melyet az információbiztonság fokozásával lehet mérsékelni. Az információbiztonsági audit és ahhoz köthető csoportok feladata, hogy értékelje a szervezetek belső kontrollkörnyezetét, mely egy komplex feladat.

A dolgozatban bemutattam az információbiztonsági kontrollok auditálásának minőségbiztosítására vonatkozó eljárásokat, melyek a mesterséges intelligencia fogalomkörébe illeszthető matematikai apparátusok révén alkalmasnak bizonyulnak az üzleti problémaként azonosított feladat közelítő megoldására (döntéstámogató rendszer – robot-auditor). A kontrollhiányosságok között meghúzódó összefüggések matematikai leképezése objektív megerősítést nyert, így a kontrollok együttes interakciói és az egyes hiányosságok konstellációi valós tudás kiaknázását teszik lehetővé, mely segítséget nyújthat a szakmai területek pontosabb és kockázatelemzés-közeleli munkavégzéséhez. Beigazolódtott, hogy a kontrollhiányosságok detektálása lehetséges felügyelt, felügyelet nélküli, valamint a két megközelítés hibrid felhasználásával, mely a kísérletekben bemutatott jóságmetrikák növelésével járt.

A gépi tanulás egyik legnagyobb kihívása az alkalmazott modellek általánosító képességének növelése, mely azt jelenti, hogy egy gépi tanuló eljárásnak még előre nem feldolgozott adathalmazon is szükséges elfogadható teljesítményt produkálnia ahhoz, hogy azt éles üzemi-környezetben fel lehessen használni. Független tesztelési adathalmaz elkülönítése indokolt az objektív mérésre, azonban a tanulórendszerrel értékes adatok megfosztásra révén lehetséges további tudás felfedezésének kockázatával kell számolni.

Megmutattam saját kereső eljárás fejlesztése és a terepmunka során gyűjtött adatvagyon feldolgozása, illetve, az adatvagyonon történő tesztelése révén, hogy a tanulóhalmaz redukálása által visszamérhetően ideálisabb eredményt lehet elérni, mely azon minták és rekordok szortírozását veszi alapul, melyek a legkevesebb hozzáadott értékkel rendelkeznek a mindenkori célváltozó becslésében. Ez azt jelenti, hogy az adattisztító és transzformáló eljárások nem kizárólag egy adathalmaz attribútumaira vonatkozhatnak, hanem rekordszinten is értelmezhetők automatizáltan, az emberi intuíció kiváltására és támogatására.

A dolgozatban szemléltettem, hogy az algoritmusok logikai kapcsolatainak és öngerjesztő mechanizmusainak révén van lehetőség a modelljóságok becslésére, mely képes orvosolni a teszhalmazok elkülönítéséből eredő tudás nem felhasználásnak kockázatát. A klasszikus értelemben vett tesztelési eljárások nélküli modell-preferencia levezetés ismertetésével tanúbizonyosságot nyerhetett az olvasó, hogy az „Univerzumot” felépítő részecskék konzisztens feldolgozásával lehetséges az objektív bizonyosságszerzés a jó kiválasztásának aspektusából, mely mindig a döntéshozó szempontja szerint értékelendő, a gép egyetlen feladata, hogy leírja a valóságot a rendelkezésre álló adatok alapján, elfogulatlanul.

8. SUMMARY

Thanks to the development of information technology, business processes are being automated to an unprecedented extent, which, in addition to modernizing repetitive manual tasks, aims to objectify and make business processes more efficient that are rich in human errors and mistakes. IT development projects abound in new types of risks that can be mitigated by enhancing the information security control environment. The objective of information security audit and related groups is to assess the internal control environment of organizations, which is a complex task.

In the dissertation, I have presented procedures to improve the quality assurance of auditing information security controls, which were proven to be suitable for the approximate solution of the task identified as the business problem (decision support system - robot auditor) through mathematical apparatus that suits to the concept of artificial intelligence. The mathematical mapping of the interrelations among control deficiencies has been objectively confirmed, so that the joint interactions of controls and the constellations of individual deficiencies allow the exploitation of real knowledge, which contributes to a more accurate and risk-based approach in professional fields. It has been demonstrated that control deficiencies can be detected using supervised and unsupervised learning, and a hybrid of the two approaches, which has increased the performance metrics presented in the experiments.

One of the biggest challenges in machine learning is to increase the generalization power of the applied models, which means that a machine learning application needs to produce acceptable performance even on previously unseen data in order to be considered acceptable in production environments. The division of the available data set into training and independent test data is justified for objective measurement, however, there is a risk of depriving valuable data and thus, discovering possible additional knowledge cannot be realized.

I have shown that more ideal results can be achieved by developing a search algorithm that aims to reduce the learning set that is based on sorting the samples and records with the least added value to the target variable. I have tested the algorithm on the collected data set and independently validated it. This means that not only to the attributes of the data set can data cleansing and transformation procedures be applied, but these procedures can also be implemented automatically at record level to elicit and support human intuition.

I have also presented that through the logical relationships and self-reinforcing mechanisms of algorithms, it is possible to estimate model performance that can remedy the risk of not utilizing the knowledge in the modelling process resulting from the divination of the data set. By deducing the model preference without testing procedures in the classical sense of the machine learning world, it has been proven that consistent processing of the particles that make up the “Universe” makes it possible to obtain objective assurance from the aspect of selecting what deems to be good, which is always to be evaluated from the point of view of the decision maker. The only task of the machine is to describe reality on the basis of the available data, impartially.

9. IRODALOMJEGYZÉK

1. ALKASASSBEH, M. (2018): A Novel Hybrid Method for Network Anomaly Detection Based on Traffic Prediction and Change Point Detection. In: *Journal of Computer Science*, 14 (2) 153-162. p.
2. ARAÚJO, B. A. (2016): Semantic Information and Artificial Intelligence. 129-140. p. In: MÜLLER V. C. (Szerk.): *Fundamental Issues of Artificial Intelligence*. Svájc: Springer International, 572 p.
3. ÁGOSTON, M. – SZLUKA, E. (1989): Tudni vagy nem tudni! Gondolatok egy nemzeti információpolitikához. Budapest: Műszaki Könyvkiadó. 178 p.
4. BAESENS, B. – VLASSELAER, V. V. – VERBEKE W. (2015): Fraud Analytics. Using Descriptive, Predictive and Social Network Techniques. A Guide to Data Science for Fraud Detection. New Jersey: Wiley. 367 p.
5. BARTA, G. (2017): A Mesterséges Intelligencia hatása az üzleti folyamatokra. In: *Magyar Internetes Agrárinformatikai Újság*, 20 (233) 1-15. p.
6. BARTA, G. (2018a): Artificial Intelligence: Blessing or Curse? In: BUSINESS AND MANAGEMENT SCIENCES: NEW CHALLENGES IN THEORY AND PRACTICE (2018)(Gödöllő). *Proceedings of the International Conference "Business and Management Sciences: New Challenges in Theory and Practice"*. Volume I. Gödöllő, p. 141-145.
7. BARTA, G. (2018b): Challenges in the compliance with the General Data Protection Regulation: Anonymization of personal information and related information security concerns. In: INTERNATIONAL SCIENTIFIC CONFERENCES OF THE FACULTY OF MANAGEMENT (10.)(2018)(Krakkó). Knowledge – Economy – Society. Business, Finance and Technology as Protection and Support for Society: proceedings. Krakkó, Cracow University of Economics. p. 115-121.
8. BARTA, G. (2018c): Implementing and Evaluating Different Machine Learning Algorithms to Predict User Localization by the Strength of User Devices' Wi-Fi Signal. In: *SEFBIS Journal*, (12) 2-11. p.
9. BARTA, G. (2018d): Predicting Human Resource Attrition with Artificial Neural Networks. 55-66. p. In: ALMÁDI B. – GARAI-FODOR M. – SZEMERE P. T. (Szerk.): *Business as usual: Comparative socio-economic studies*. Budapest: Vízkapu Kiadó, 127 p.
10. BARTA, G. (2018e): The Increasing Role of IT Auditors in Financial Audit: Risks and Intelligent Answers. In: *Business Management and Education*, 16 (1) 81-93. p.
11. BARTA, G. (2020): Tanúsítványok értékelése ellátási láncok IT biztonsági megfelelésének vizsgálatára. In: *Logisztikai trendek és legjobb gyakorlatok*, 6 (1) 27-30. p.
12. BARTA, G. – GÖRCSEI, G. (2017): Intelligent Decision Making and Process Automation for Public Organizations. In: INTERNATIONAL SCIENTIFIC CORRESPONDENCE CONFERENCE (5.)(2017)(Nitra). Legal, economic, managerial and environmental aspects of performance competencies by local authorities. Nitra, Slovak University of Agriculture in Nitra. p. 30-37.
13. BARTA, G. – GÖRCSEI, G. (2018): Artificial Intelligence and Audit: Why is it necessary to audit the intelligent decision support? In: KÖZGAZDÁSZ DOKTORANDUSZOK ÉS KUTATÓK TÉLI KONFERENCIÁJA (4.)(2018)(Gödöllő). Közgazdász Doktoranduszok és Kutatók IV. Téli Konferenciája: Konferenciakötet. Budapest: Doktoranduszok Országos Szövetsége. p. 225-234.
14. BARTA, G. – GÖRCSEI, G. (2019): Csevegőrobotok a vállalati működésben. In: GAZDÁLKODÁS ÉS MENEDZSMENT TUDOMÁNYOS KONFERENCIA (3.)(2019)(Kecskemét). Versenyképesség és innováció. Kecskemét, Neumann János Egyetem Kertészeti és Vidékfejlesztési Kar. p. 912-917.

15. BARTA, G. – GÖRCSI, G. (2020): Assessing and managing business risks for artificial intelligence based business process automation. In: SCIENTIFIC CONFERENCE ON CONTEMPORARY ISSUES IN BUSINESS, MANAGEMENT AND ECONOMIC ENGINEERING (6.)(2019)(Vilnius). 6th International Scientific Conference Contemporary Issues in Business, Management and Economics Engineering '2019: proceedings. Vilnius, Vilnius Gediminas Technical University Press. p. 823-832.
16. BARTA, G. – GÖRCSI, G. (2021): Risk Management Considerations for Artificial Intelligence Business Applications. In: *International Journal of Economics and Business research*, 21 (1) 87-106. p.
17. BARTA, G. – ŁĘTEK, M. (2015): Non-financial Performance Indicators as a New Approach to Measure the Value of Companies. In: INTERNATIONAL SCIENTIFIC CONFERENCES OF THE FACULTY OF MANAGEMENT (7.)(2015)(Krakkó). Knowledge-Economy-Society: Reorientation of paradigms and concepts of management in the contemporary economy. Krakko, Cracow University of Economics. 247-254. p.
18. BARTA, G. – LUDVAI, N. – PUSKÁS, A. (2020): The analysis of data privacy incidents and sanctions in Europe after GDPR enforcement. In: INTERNATIONAL WINTER CONFERENCE OF ECONOMICS PHD STUDENTS AND RESEARCHERS (6.)(2020)(Gödöllő). VI. International Winter Conference of Economics PhD Students and Researchers: Conference Proceedings. Budapest, Association of Hungarian PhD and DLA Students. p. 35-48.
19. BARTA, G. – PITLIK, L. (2018a): A Titanic katasztrófa túlélőinek becslése döntési fa alapú gépi tanuló eljárással. In: *Magyar Internetes Agrárinformatikai Újság*, 21 (234) 1-24. p.
20. BARTA, G. – PITLIK, L. (2018b): Startup felvásárlások multikulturális hátterének elemzése, avagy mesterséges intelligencia alapú ellenőrzőszámítás diszkriminancia-elemzéshez. 15-37. p. In: FARKAS A. (Szerk.): *A gazdaság kulturális szerkezete*. Gödöllő: Szent István Egyetemi Kiadó, 240 p.
21. BARTA, G. – PITLIK, L. (2020): Hipotézis-tervezés PhD-disszertációkhoz - Konzisztens gépi tanuló modellezés beltéri felhasználói lokalizáció meghatározásának pontosítására. In: *Magyar Internetes Agrárinformatikai Újság*, 23 (263) 1-13. p.
22. BÁNKUTI, GY. (2010): About the method of component-based object comparison for objectivity. In: INTERNATIONAL CONGRESS OF MATHEMATICIANS. (2010)(Hindustan). International Congress of Mathematicians: Abstracts, short communications, posters. Hindustan, Book Agency. p. 593-594.
23. BEALUC, C. – ROSENTHAL, J. S. (2018): Handling Missing Values using Decision Trees with Branch-Exclusive Splits. <https://arxiv.org/pdf/1804.10168.pdf>. Keresőprogram: Google. Kulcsszavak: handling missing values using machine learning. Lekérdezés időpontja: 2020.01.17.
24. BELLMAN, R. (1978): An introduction to artificial intelligence: can computers think? San Francisco: Boyd & Fraser Publishing Company. 146 p. Idézve: RUSSEL S. – NORVIG P. (2005): Mesterséges Intelligencia modern megközelítésben. Budapest: Panem Kiadó. 1206 p.
25. BHAGOJI, A. N. et al. (2018): Enhancing Robustness of Machine Learning Systems via Data Transformation. In: ANNUAL CONFERENCE ON INFORMATION SCIENCES AND SYSTEMS (52.)(2018)(Princeton). 52nd Annual Conference on Information Sciences and Systems (CISS): proceedings. Princeton, p. 1-5.
26. BHUYAN, M. H. – BHATTACHARYYA, D. K. – KALITA, J. K. (2011): Survey on Incremental Approaches for Network Anomaly Detection. In: *International Journal of Communication Networks and Information Security (IJCNIS)*, 3 (3) 226-239. p.

27. BIRÓ, B. – PITLIK, L. (2020): The evaluation of exogenous ligands cross-reactivity to a 7TM receptor based on online artificial intelligence engines. In: *Magyar Internetes Agrárinformatikai Újság*, 23 (257) 1-35. p.
28. BODA, GY. – JUHÁSZ, P. – STOCKER, M. (2009): A tudás mint termelési tényező. In: *Közgazdaság*, 4 (3) 117-132. p.
29. BODA, M. A. (2019): Üzleti szimulációk és tanuló-rendszerek döntéshozatali mechanizmusai. Doktori disszertáció. Gödöllő: Szent István Egyetem, Gazdálkodás és Szervezéstudományok Doktori Iskola. 214 p.
30. BODON, F. (2010): Adatbányászati algoritmusok. Budapest: Free Software Foundation. 278 p.
31. BORGULYA, I. (1998): Neurális hálók és fuzzy-rendszerek. Budapest – Pécs: Dialóg Campus Szakkönyvek. 226 p.
32. BUNKÓCZI, L. (1998): Mesterséges intelligencia alapú prognosztikai modulok adaptálása a EU/SPEL-Hungary rendszerhez az alapadatbázisok konzisztenciájának egyidejű ellenőrzésével. In: *Magyar Internetes Agrárinformatikai Újság*. <https://miau.myx.hu/miau/04/blaci.html> Keresőprogram: Google. Kulcsszavak: mesterséges intelligencia definiálása. Lekérdezés időpontja: 2020. 11. 27.
33. BURKOV, A. (2019): The Hundred-Page Machine Learning Book. Független kiadó. 160 p.
34. BUXBAUM, M. (2006): Vállalati internal audit a gyakorlatban. Hogyan válaszol a belső ellenőrzés a vállalkozói szféra kihívásaira? Budapest: ETK Szolgáltató. 143 p.
35. BÜCKER, M. et al. (2020): Transparency, Auditability and eXplainability of Machine Learning Models in Credit Scoring. <https://arxiv.org/pdf/2009.13384.pdf> Keresőprogram: Google. Kulcsszavak: Machine learning and audit. Lekérdezés időpontja: 2020. 11. 17.
36. CAPURRO, R. (1992): What is Information Science for? A philosophical reflection. In: CONCEPTION OF LIBRARY AND INFORMATION SCIENCE. (1992)(TAMPERE). Conception of Library and Information Science: proceedings. Tampere, p. 82-98.
37. CERDA, P. – VAROQUAUX, G. – KÉGL, B. (2018): Similarity encoding for learning with dirty categorical variables. In: *Mach Learn*, (107) 1477–1494. p.
38. CHARNIAK, E. – MCDERMOTT, D. (1985): Introduction to Artificial Intelligence. Massachusetts: Addison-Wesley. 701 p. Idézve: RUSSEL S. – NORVIG P. (2005): Mesterséges Intelligencia modern megközelítésben. Budapest: Panem Kiadó. 1206 p.
39. CHIKÁN, A. (2008): Vállalatgazdaságtan. 4. átdolgozott, bővített kiadás. Budapest: Aula Kiadó. 616 p.
40. CHOLLET, F. (2018): Deep Learning with Python. New York: Manning. 361 p.
41. COMTE, A. (2009): The positive philosophy of Auguste Comte. Digitálisan újranyomtatott verzió az eredeti 1853-an kiadott verzió alapján. Cambridge: Cambridge University Press. 524 p.
42. COSO (2020): About us. <https://www.coso.org/Pages/aboutus.aspx> Keresőprogram: Google. Kulcsszavak: Coso framework. Lekérdezés időpontja: 2020. 11. 27.
43. CURTIS, P. – CAREY, M. (2012): Risk assessment in practice. Deloitte & Touche LLP. 28 p.
44. DAI, J. – FAZELPOUR, S. – LIPTON, Z. C. (2020): Fair Machine Learning Under Partial Compliance. <https://arxiv.org/pdf/2011.03654.pdf> Keresőprogram: Google. Kulcsszavak: fair machine learning. Lekérdezés időpontja: 2020. 11. 17.
45. DANGETI, P. (2017): Statistics for Machine Learning. Birmingham: Packt. 425 p.
46. DARÓCZY, B. – FRIEDL, K. – KABÓDI, L. – PERESZLÉNYI, A. – SZABÓ, D. (2021): Quantum Inspired Adaptive Boosting. <https://arxiv.org/pdf/2102.00949.pdf> Keresőprogram: Google. Kulcsszavak: adaboost. Lekérdezés időpontja: 2021. 05. 12.

47. DAVIS, C. – SCHILLER, M. – WHEELER, K. (2011): IT Auditing: Using Controls to Protect Information Assets. Second Edition. USA: McGraw-Hill. 480 p.
48. DEAN, J. (2020): The Deep Learning Revolution and Its Implications for Computer Architecture and Chip Design. In: IEEE INTERNATIONAL CONFERENCE ON SOLID-STATE CIRCUITS (2020)(San Francisco). IEEE International Conference on Solid-State Circuits (ISSCC): proceedings. San Francisco, p. 8-14.
49. DELGADO, M. F. – CERNADAS, E. – BARRO, S. – AMORIM, D. (2014): Do we Need Hundreds of Classifiers to Solve Real World Classification Problems? In: *Journal of Machine Learning Research*, (15) 3133-3181. p.
50. DENNING, D. (1987): An intrusion-Detection Model. In: IEEE TRANSACTION ON SOFTWARE ENGINEERING (1986)(Oakland). IEEE Transactions on Software Engineering: proceedings. Oakland, p. 118-131.
51. DIAMANT, E. (2017): Advances in Artificial Intelligence: Are you sure, we are on the right track? In: *Transactions on networks and communications*, 5 (4) 23-30. p.
52. DILEK, S. et al. (2015): Applications of artificial intelligence techniques to combating cyber crimes: a review. In: *International Journal of Artificial Intelligence & Applications (IJAIA)*, 6 (1) 21-39. p.
53. DOBAY, P. (1997): Vállalati információmenedzsment. Budapest: Nemzeti Tankönyvkiadó. 310 p.
54. DOBÓ, A. (1992): A hasonlóságelmélet alkalmazása a Joker rendszerben. Budapest: Prodinform. 62 p.
55. DUA, S. – DU, X. (2011): Data Mining and Machine Learning in Cybersecurity. USA: Taylor and Francis. 223 p.
56. ESKIN, E. – ARNOLD, A. – PRERAU, M. (2002): A geometric framework for unsupervised anomaly detection: Detecting intrusions in unlabeled data. In: *Applications of Data Mining in Computer Security*, (6) 1-20. p.
57. Európai Elméleti Számítástudományi Egyesület (2003): Gödel Prize – 2003. <https://eatcs.org/index.php/component/content/article/505> Keresőprogram: Google. Kulcsszavak: Gödel Prize 2003. Lekérdezés időpontja: 2020. 06. 17.
58. FANAEE, T. H. – OLIVEIRA, M. – GAMA, J. (2018): Event and Anomaly Detection Using Tucker3 Decomposition. <https://arxiv.org/pdf/1406.3266.pdf>. Keresőprogram: Google. Kulcsszavak: event and anomaly detection. Lekérdezés időpontja: 2020. 02. 22.
59. FARKASNÉ, F. M. – MOLNÁR, J. (2017): Közgazdaságtan I. Mikroökonómia. Debrecen: Debreceni Egyetem Agrár- és Műszaki Tudományok centruma. Agrárgazdasági és Vidékfejlesztési Kar. 274 p.
60. FEINSTEIN, L. – SCHNACKENBERG, D. – BALUPARI, R. – KINDRED, D. (2003): Statistical approaches to Ddos attack detection and response. In: DARPA INFORMATION SURVIVABILITY AND EXPOSITION (2003)(Washington). DARPA Information Survivability Conference and Exposition: proceedings. Washington, p. 303-314.
61. FORGÓ, S. (2011): A kommunikációelmélet alapjai. Eger: Médiainformatikai Kiadványok. 172 p.
62. FREUND, Y. – SCHAPIRE, E. R. (1996): Experiments with a new boosting algorithm. In: MACHINE LEARNING INTERNATIONAL CONFERENCE (13.)(1996)(San Diego). Machine Learning: Proceedings of the Thirteenth International Conference. San Diego, p. 1-9.
63. FRIEDMAN, J. H. (2001): Greedy function approximation: A gradient boosting machine. In: *The Annals of Statistics*, 29 (5) 1189-1232. p.

64. FÜLÖP, G. (1996): Az információ. 2. bővített és átdolgozott kiadás. Budapest: Eötvös Loránd Tudományegyetem, Könyvtártudományi – Informatikai Tanszék kiadványa. 236 p.
65. GELBOWITZ, A. (2021): Decision Trees and Random Forests Guide: An Overview Of Decision Trees And Random Forests: Machine Learning Design Patterns. Független Kiadó. 65 p.
66. GÉRON, A. (2017): Hands-On Machine Learning with Scikit-Learn & Tensorflow. Concepts, tools, and technologies to build intelligent systems. USA: O'Reilly. 549 p.
67. GHOSH, A. K. – WANKEN, J. – CHARRON, F. (1998): Detecting anomalous and unknown intrusions against programs. In: Computer Security Applications Conference (1998)(Phoenix). Proceedings of the 1998 Annual Computer Security Applications Conference (ACSAC). Phoenix, p. 1-9.
68. GOODFELLOW, I. J. – POUGET-ABADIE, J. – MIRRZA, M. – XU, B. – WARDEFARLEY, D. – OZAIR, S. – COURVILLE, A. – BENGIO, Y. (2014): Generative Adversarial Networks. In: *Advances in Neural Information Processing Systems*, 3 (11) 1-9. p.
69. GOODFELLOW, I. A. – BENGIO, Y. – COURVILLE, A. (2016): Deep Learning. USA: MIT Press. 775 p.
70. GOOGLE BOOKS NGRAM VIEWER (2020): https://books.google.com/ngrams/graph?content=artificial+intelligence%2Cbig+data%2Cdata+mining%2COLAP%2Cmachine+learning&year_start=1950&year_end=2019&corpus=26&smoothing=3 Lekérdezés időpontja: 2020. 09. 27.
71. GÖRCSI, G. – BARTA, G. (2018): A CRM rendszerek szerepe a vevőkapcsolatok stratégiai kezelésében, vevőszegmentációs döntésekben. In: In: KÖZGAZDÁSZ DOKTORANDUSZOK ÉS KUTATÓK TÉLI KONFERENCIÁJA (4.)(2018)(Gödöllő). Közgazdász Doktoranduszok és Kutatók IV. Téli Konferenciája: Konferenciakötet. Budapest: Doktoranduszok Országos Szövetsége. p. 26-33.
72. GÖRCSI, G. – BARTA, G. (2019): Az információs rendszer szerepe a döntési folyamatban. In: GAZDÁLKODÁS ÉS MENEDZSMENT TUDOMÁNYOS KONFERENCIA (3.)(2019)(Kecskemét). Versenyképesség és innováció. Kecskemét, Neumann János Egyetem Kertészeti és Vidékfejlesztési Kar. p. 252-256.
73. GÖRCSI, G. – SZÉLES, ZS. – BARTA, G. (2019): Üzleti intelligencia megoldások alkalmazásának sikertényezői - A hazai szolgáltató szektor nagyvállalatainak körében végzett mélyinterjú kutatás. In: *Információs Társadalom: Társadalomtudományi Folyóirat*, 19 (2) 23-34. p.
74. GREENHOUSE L. (2005): Justices Unanimously Overturn Conviction of Arthur Andersen. <https://www.nytimes.com/2005/05/31/business/justices-unanimously-overturn-conviction-of-arthur-andersen.html> Keresőprogram: Google. Kulcsszavak: Arthur Andersen. Lekérdezés időpontja: 2020. 08. 08.
75. GUIDOTTI, R. – MONREALE, A. – RUGGIERI, S. – TURINI, F. – GIANNOTTI, F. – PEDRESCHI, D. (2018): A Survey of Methods for Explaining Black Box Models. In: *ACM Computing Surveys*, 51 (5) 1-42. p.
76. HACKELING, G. (2014): Mastering Machine Learning with scikit-learn. Birmingham: Packt. 238 p.
77. HADJERES, G. – PACHET, F. – NIELSEN, F. (2017): DeepBach: a steerable model for bach chorales generation. In: INTERNATIONAL CONFERENCE ON MACHINE LEARNING (70.)(2017)(Sydney). ICML'17: Proceedings of the 34th International Conference on Machine Learning. Sydney, 1362-1371. p.
78. HAES S. D. et al. (2018): COBIT 2019 Framework: Governance and Management Objectives. Schaumburg: ISACA. 302 p.

79. HAN, J. – KAMBER, M. (2011): Data Mining: Concepts and Techniques. 3rd Edition. USA: Morgan Kaufmann. 744 p.
80. HARKEVICS, A. A. (1960): O cennoszti informacii. Problemu kibernetiki. Vűp. 4. Moszkva. Idézi: FÜLÖP, G. (1996): Az információ. 2. bővített és átdolgozott kiadás. Budapest: Eötvös Loránd Tudományegyetem, Könyvtártudományi – Informatikai Tanszék kiadványa. 236 p.
81. HASTIE T. – TIBSHIRANI, R. – FRIEDMAN, J. (2009): The Elements of Statistical Learning. Data Mining, Inference, and Prediction. Second Edition. New York: Springer. 745 p.
82. HAUGELAND, J. (1985): Artificial Intelligence: The Very Idea. Massachusetts: MIT Press, 299 p. Idézve: RUSSEL S. – NORVIG P. (2005): Mesterséges Intelligencia modern megközelítésben. Budapest: Panem Kiadó. 1206 p.
83. HAWKING, S. W. (2017): Fekete lyukak. Budapest: Akkord Kiadó. 116 p.
84. HEARTY, J. (2016): Advanced Machine Learning with Python. Birmingham: Packt. 254 p.
85. HOLZINGER, A. (2018): Explainable AI (ex-AI). Graz: Holzinger. 6 p.
86. HORVÁTH, G. (2006): Neurális hálózatok és műszaki alkalmazásaik. Budapest: Műszaki Egyetem Kiadó. 314 p.
87. HORVÁTH, Z. – JENAK, I. – BRACHMANN, F. (2016): Battery consumption of smartphone sensors. In: *Journal of Reliable Intelligent Environments*, 3 (2) 131-136. p.
88. Information Security Forum (2014): IRAM2. The next generation of assessing information risk. Information Security Forum Limited. 103 p.
89. Information System Audit & Control Association (2015): CISA Review Manual 2015. USA: ISACA. 426 p.
90. Information System Audit & Control Association (2016): CISM Review Manual 14th edition. USA: ISACA. 283 p.
91. Internet World Stats (2020): World Internet Usage and Population Statistics. 2020 Year-Q4 Estimates. <https://www.internetworldstats.com/stats.htm> Keresőprogram: Google. Kulcsszavak: Internet world statistics. Lekérdezés időpontja: 2020. 11. 19.
92. ISO (2018): Te ISO Survey of Management System Standard Certifications 2018. <https://isotc.iso.org/livelink/livelink?func=ll&objId=18808772&objAction=browse&viewType=1> Keresőprogram: Google. Kulcsszavak: ISO survey. Lekérdezés időpontja: 2020. 09. 22.
93. IZZA, Y. – IGNATIEV, A. – MARQUES-SILVA, J. (2020): On Explaining Decision Trees. <https://arxiv.org/pdf/2010.11034.pdf> Keresőprogram: Google. Kulcsszavak: Decision trees. Lekérdezés időpontja: 2021. 05. 12.
94. KÁSA, R. (2011): Neurális fuzzy rendszerek alkalmazása a társadalomtudományi kutatásban az innovációs potenciál mérésére. Doktori disszertáció. Miskolc: Miskolci Egyetem, Gazdaságtudományi Kar, Vezetéstudományi Intézet. 252 p.
95. KÁSA, R. (2018): Neurális hálók alkalmazásának lehetőségei innovációs teljesítmény mérésére. In: *LOGISZTIKA - INFORMATIKA – MENEDZSMENT*, 3 (1) 60-73. p.
96. KISS, T. J. (2016): A tudásgazdaság jellemzői Magyarország vonatkozásában. In: *International Journal of Engineering and Management Sciences (IJEMS)*, 1 (1) 1-11. p.
97. KLEIN, A. (2017): Hardver Drive Cost Per Gigabyte. <https://www.backblaze.com/blog/hard-drive-cost-per-gigabyte/> Keresőprogram: Google. Kulcsszavak: Hard drive cost per gigabyte. Lekérdezés időpontja: 2020. 10. 24.
98. KNUTH, D. (1995): A=B. Előszó PETKOVSEK, M. – WILF, S. H. – ZEILBERGER, D. könyvében. Massachusetts: A K Peters/CRC Press. 224 p.
99. KOMENCZI, B. (2011). Információelmélet. Eger: Médiainformatikai Kiadványok. 131 p.

100. KOMOROWSKI, M. (2014): A history of storage cost (update). <https://mkomo.com/cost-per-gigabyte-update> Keresőprogram: Google. Kulcsszavak: History of storage cost. Lekérdezés időpontja: 2020. 10. 24.
101. KOVÁCS, E. (2014): Többváltozós adatelemzés. Budapest: Typotex Kiadó. 252 p.
102. KOVÁCS, Z. (2009): Szimulációs eszközök és megoldások műszaki és gazdasági rendszerekben. In: INNOVÁCIÓ AZ EGYETEMI KÉPZÉSBEN ÉS KUTATÁSBAN (2.)(2009)(Balatonvilágos). „Innováció az egyetemi képzésben és kutatásban” Jubileumi Tudományos Konferencia. Balatonvilágos, előadásjegyzet.
103. KURZWEIL, R. (1990): The age of intelligent machines. Massachusetss: MIT Press. 565 p.
104. LANGEFORS, B. (1973): Theoretical analysis of information systems. Auerbach: Wiley. 502 p.
105. LAPAN, M. (2018): Deep Reinforcement Learning Hands-On. Birmingham: Packt. 523 p.
106. LAUDON, K. C. – LAUDON, J. P. (1991): Management Information Systems. A Contemporary Perspective. USA: Macmillan. 336 p.
107. LAUDON, K. C. – LAUDON, J. P. (2015): Management Information Systems. Global Edition. USA: Pearson. 648 p.
108. LEE, W. – STOLFO, S. J. (2000): A framework for constructing features and models for intrusion detection systems. In: *ACM Transactions on Information Systems Security*, 3 (4) 227-260. p.
109. LENCSES, E. – DUNAY, A. – MÉSZÁROS, K. – KOVÁCS, A. (2019): Fejőrobot technológia bevezetésének hatása az állategészségügyi költségekre. In: GÖDÖLLŐI ÁLLATTENYÉSZTÉSI TUDOMÁNYOS NAP (7.)(2019)(Gödöllő). Előadások és posztterek összefoglaló kötete. Gödöllő, Szent István Egyetemi Kiadó. p. 26-26.
110. LEUNG, K. – LECKIE, C. (2005): Unsupervised anomaly detection in network intrusion detection using clusters. In: AUSTRALASIAN CONFERENCE ON COMPUTER SCIENCE (38.)(2005)(Newcastle). ACSC '05 Proceedings of the Twenty-eighth Australasian conference on Computer Science. Newcastle, p. 333-342.
111. LIAO, Y.H. – VEMURI, V. R. (2002): Use of k-nearest neighbor classifier for intrusion detection. In: *Computer & Security*, 21 (5) 439-448. p.
112. LIN, Y. H. – BRADY J. P. – FORMAN-KAY, J. D. – CHAN, H. S. (2017): Charge Pattern Matching as a “Fuzzy” Mode of Molecular Recognition for the Functional Phase Separations of Intrinsically Disordered Proteins. In: *New Journal of Physics*, 19 (11) 1-23 p.
113. LIU, F. – SHI, Y. – LI, P. (2017): Analysis of the Relation between Artificial Intelligence and the Internet from the Perspective of Brain Science. In: *Procedia Computer Science*, (122) 377-383. p.
114. MAGDA R. (2015): The effects of globalisation on logistics in Europe and in Hungary. In: *Logistics and Transport*, (26) 33-42. p.
115. Magyar Elektronikus Könyvtár (2016): A Magyar Értelmező Nyelv Szótára. <https://mek.oszk.hu/> Keresőprogram: Google. Kulcsszavak: magyar nyelv értelmező szótára gyanú. Lekérdezés időpontja: 2020. 08. 17.
116. Magyar Könyvvizsgálói Kamara (2015): Brókerbotrány: a könyvvizsgálók az Alkotmánybírósághoz fordulhatnak. https://mkvk.hu/hu/kamarai/kozlemenyek/sajtotajekoztato_20150402 Keresőprogram: Google. Kulcsszavak: Buda-Cash Brókerbotrány, könyvvizsgálat. Lekérdezés időpontja: 2020. 08. 02.
117. Magyar Nemzeti Bank (2015): A Jegybank azonnali hatállyal felfüggesztette a Buda-Cash Brókerház működési engedélyét és felügyeleti biztosokat rendelt ki. <https://www.mnb.hu/sajtoszoba/sajtokozlemenyek/2015-evi-sajtokozlemenyek/a-jegybank-azonnali-hatallyal-felfuggesztette-a-buda-cash-brokerhaz-mukodesi-engedelyet-es-felugy>

életi-biztosokat-rendelt-ki Keresőprogram: Google. Kulcsszavak: MNB sajtószoja, Buda-Cash Brókerház. Lekérdezés időpontja: 2020. 05. 02.

118. Magyar Tudományos Akadémia (2017): Tudományági nomenklatúra. <https://mta.hu/doktoritanacs/tudomanyagi-nomenklatura-106809> Keresőprogram: Google. Kulcsszavak: Tudományági nomenklatúra. Lekérdezés időpontja: 2020. 11. 19.
119. MAHONEY, M. V. – CHAN, P. K. (2003): Learning Rules for Anomaly Detection of Hostile Network Traffic. In: IEEE INTERNATIONAL CONFERENCE ON DATA MINING (3.)(2003)(Melbourne). ICDM '03 Proceedings of the Third IEEE International Conference on Data Mining. Melbourne, p. 601-611.
120. MAQUEDA, A. I. – LOQUERCIO, A. – GALLEGO, G. – GARCÍA, N. – SCARAMUZZA, D. (2018): Event-based Vision meets Deep Learning on Steering Prediction for Self-driving Cars. In: CONFERENCE ON COMPUTER VISION AND PATTERN RECOGNITION (2018)(Salt Lake City). IEEE/CVF Conference on Computer Vision and Pattern Recognition: proceedings. Salt Lake City, p. 5419-5427.
121. MASON, L. – BAXTER, J. – BARTLETT, P. – FREAN, M. (1999): Boosting algorithms as gradient descent. In: INTERNATIONAL CONFERENCE ON NEURAL INFORMATION PROCESSING SYSTEMS (12.)(1999)(Cambridge). NIPS'99: Proceedings of the 12th International Conference on Neural Information Processing Systems. Cambridge, p. 512-518.
122. MATA, J. – MIGUEL, I. – DURÁN, R. J. – MERAYO, N. – SINGH, S. K. – JUKAN, A. – CHAMANIA, M. (2018): Artificial intelligence (AI) methods in optical networks: A comprehensive survey. In: *Optical Switching and Networking*, (28) 43-57. p.
123. MÁTYUS, I. (2015): Tudományos tudás az információs társadalomban. Oktatási segédanyag. TÁMOP-4.2.1.D-15/1/KONYV-2015-0002 azonosítószámú pályázat keretében készült. 17 p.
124. MCCARTHY J. – MINSKY, M. L. – ROCHESTER, N. – SHANNON, C. E. (1955): Proposal for the Dartmouth Summer Research Project on Artificial Intelligence. In: *AI Magazine*, 27 (4) 12-14. p.
125. MCCLURE, N. (2017): TensorFlow Machine Learning Cookbook. Birmingham: Packt. 351 p.
126. MCCULLOCH W. S. – PITTS W. (1943): A Logical Calculus of the Ideas Immanent in Nervous Activity. In: *The bulletin of mathematical biophysics*, 5 (4) 115-133. p.
127. MEASE, D. – WYNER, A. (2008): Evidence Contrary to the Statistical View of Boosting. In: *Journal of Machine Learning Research*, (9) 131-201. p.
128. Mesterséges Intelligencia Koalíció (2020): Magyarország Mesterséges Intelligencia Stratégiája. Budapest: Digitális Jólét, 60 p.
129. MÉRŐ, L. (2007): Mindenki másképp egyforma. A játékelmélet és a racionalitás pszichológiája. Budapest: Tericum Kiadó. 392 p.
130. MOLNÁR, B. – KŐ, A. (2009): Információrendszerek auditálása. Az informatika és az információrendszerek ellenőrzési és irányítási módszerei. Budapest: Corvinho Technology Transfer. 388 p.
131. MUNK, S. (2007): Katonai informatikai a XXI. század elején. Budapest: Zrínyi Kiadó. 264 p.
132. NG, A. (2018): Machine Learning Yearning. Technical Strategy for AI Engineers, In the Era of Deep Learning. DeepLearning.ai. 118 p.
133. NILLSON, N. J. (1998): Artificial Intelligence: A new Synthesis. San Mateo: Morgan Kaufmann. 513 p. Idézve: RUSSEL, S. – NORVIG, P. (2005): Mesterséges Intelligencia modern megközelítésben. Budapest: Panem Kiadó. 1206 p.
134. OLÁH, J. – POPP, J. – ERDEI, E. (2019): Az Ipar 5.0 megjelenése: ember és robot együttműködése. In: *Logisztikai trendek és legjobb gyakorlatok*, 5 (1) 12-19. p.

135. PALUZO-HIDALGO, E. – GONZALEZ-DIAZ, R. – GUTIÉRREZ-NARANJO, M. A. (2020): Two-hidden-layer feed-forward networks are universal approximators: A constructive approach. In: *Neural Networks*, (131) 29-36. p.
136. PANIGUTTI, C. (2020): FairLens: Auditing Black-box Clinical Decision Support Systems. <https://arxiv.org/pdf/2011.04049.pdf> Keresőprogram: Google. Kulcsszavak: Machine learning and audit. Lekérdezés időpontja: 2020. 11. 17.
137. PETŐ, I. (2013): Hasonlóságelemzés COCO használatával. Oktatási segédanyag. https://miau.my-x.hu/miau/189/coco_demo.pdf Keresőprogram: Google. Kulcsszavak: Hasonlóságelemzés oktatási segédanyag. Lekérdezés időpontja: 2021. 05. 12.
138. PFEFFER, A. – RUTTENBERG, B. E. – KELLOGG, L. – HOWARD, M. – CALL, C. – O’CONNOR, A. – TAKATA, G. – REILLY, S. – PATTEN, T. – TAYLOR, T. – HALL, R. – LAKHOTIA, A. – MILES, C. – SCOFIELD, D. – FRANK, J. (2017): Artificial Intelligence Based Malware Analysis. <https://arxiv.org/pdf/1704.08716.pdf>. Keresőprogram: Google. Kulcsszavak: AI based malware analysis. Lekérdezés időpontja: 2020. 04. 02.
139. PITLIK, L. (2013): Gyanúgenerálás a HR-kockázatok minimalizálása érdekében – hasonlóságelemzéssel. In: *Tudásmenedzsment. A Pécsi Tudományegyetem Felnőttképzési és Emberi Erőforrás Fejlesztési Karának periodikája*, 4 (1) 171-178. p.
140. PITLIK, L. (2014): My-X team, avagy egy innovatív „ötlet-istálló”. Budapest: Innoreg Közép-magyarországi Regionális Innovációs Ügynökség. 28 p.
141. PITLIK, L. – BUNKÓCZI, L. – PETŐ, I. (2005): Environmental-Ecological Consistencies in automation of modelling. In: *HUNGARIAN BIOMETRIC AND BIOMATHEMATICS (8.)(2005)(Budapest)*. VII. Hungarian Biometric and Biomathematics Conference: proceedings. Budapest, p. 1-7.
142. PITLIK, L. – PITLIK, M. (2021a): A multikollinearitás téves kezelése, avagy információvesztést generáló változókizárás a hasonlóságelemzésben. In: *Magyar Internetes Agrárinformatikai Újság*, 24 (270) 1-9. p.
143. PITLIK, L. – PITLIK, M. (2021b): Solver-függő alternatív megoldások a hasonlóságelemzésben. In: *Magyar Internetes Agrárinformatikai Újság*, 24 (274) 1-8. p.
144. PITLIK, L. – PITLIK, M. – PITLIK, M. (2020a): How to design a graphical expert system with teaching/learning effects? In: *Magyar Internetes Agrárinformatikai Újság*, 23 (258) 1-7. p.
145. PITLIK L. – PITLIK, M. – BARTÓK, P. – RIKK, J. (2020b): Értékek matematikája. In: *Magyar Internetes Agrárinformatikai Újság*, 23 (267) 1-9. p.
146. PITLIK, L. – RIKK, J. – GÁNGÓ, V. – TÓTH, CS. (2020c): A távoktatás, mint kritikus oktatási üzem – IT-aspektusai, avagy felkészülés a duális képzésre. In: *Magyar Internetes Agrárinformatikai Újság*, 23 (266) 1-26. p.
147. PITLIK, L. – VARGA, Z. – BARTA, G. – LOSONCZI, GY. – PITLIK, L. (jun.) – PITLIK, M. – PITLIK, M. (2017): Magyar statisztikai régiók érintettségi sorrendje a szálláshelyek árbevételének havi adatai alapján eltérő módszertanokkal. In: *VIDÉKFEJLESZTÉSI KONFERENCIA (1.)(2017)(Szarvas)*. Magyar vidék - perspektívák, megoldások a XXI. században. Szarvas, Szent István Egyetem Egyetemi Kiadó. p. 127-140.
148. POMPON, R. (2016): IT Security Risk Control Management. An Audit Preparation Plan. Washington: Apress. 311 p.
149. POOLE D. – MACKWORTH, A. – GOEBEL, R. (1998): Computational Intelligence. A logical approach. New York: Oxford University Press. 576 p. Idézve: RUSSEL S. – NORVIG P. (2005): Mesterséges Intelligencia modern megközelítésben. Budapest: Panem Kiadó. 1206 p.

150. POÓR J. – SASVÁRI, P. – SZALAY, ZS. – PETŐ, I. – GYURIÁN, N. – SUHAJDA, CS. J. – ZSIGRI, F. (2020): The implementation and management of e-learning in companies – the state of e-learning in Hungary based on empirical research. In: *Journal of Engineering Management and Competitiveness*, 10 (1) 3-14. p.
151. PORTNOY, L. – ESKIN, E. – STOLFO, S. (2001): Intrusion detection with unlabeled data using clustering. In: WORKSHOP ON DATA MINING APPLIED SECURITY (2001). Proceedings of ACM CSS Workshop on Data Mining Applied Security. p. 1-25.
152. PROVOST, F. – FAWCETT, T. (2013): Data Science for Business: What You Need to Know about Data Mining and Data-Analytic Thinking. USA: O'Reilly. 384 p.
153. RANJAN, R. – SAHOO, G. (2014): A new clustering approach for anomaly intrusion detection. In: *International Journal of Data Mining & Knowledge Management Process*, 4 (2) 29-38. p.
154. RASCHKA, S. (2015): Python Machine Learning. Birmingham: Packt. 425 p.
155. RASCHKA, S. – MIRJALILI, V. (2017): Python Machine Learning - Second Edition: Machine Learning and Deep Learning with Python, scikit-learn, and TensorFlow. Birmingham: Packt. 622 p.
156. RICH E. et al. (2009): Artificial Intelligence. Third Edition. New York: McGraw-Hill. 568 p.
157. RUSSEL, S. – NORVIG, P. (2005): Mesterséges Intelligencia modern megközelítésben. Második Kiadás. Budapest: Panem Kiadó. 1206 p.
158. RUSSEL S. – NORVIG P. (2009): Artificial Intelligence: A Modern Approach. 3rd Edition. USA: Pearson. 1152 p.
159. SAGAR, G. V. R. (2015): Modeling of Artificial Neural Networks using Evolutionary Algorithms. Germany: Lambert Academic Publishing. 179 p.
160. SAJTOS, L. – MITEV, A. (2009): SPSS kutatási és adatelemzési kézikönyv. Budapest: Alinea Kiadó. 402 p.
161. SCOPUS (2021a): <https://www.scopus.com/term/analyzer.uri?sid=a041b6c92c025536929a745e3b45dc1a&origin=resultslist&src=s&s=TITLE-ABS-KEY%28artificial+intelligence%29&sort=plf-f&sdt=b&sot=b&sl=38&count=381294&analyzeResult=s=Analyze+results&txGid=a391bd72a2e5e5a2a5a61f6d262db437> Lekérdezés időpontja: 2020. 01. 01.
162. SCOPUS (2021b): [https://www2.scopus.com/results/results.uri?src=s&st1=&st2=&sot=b&sdt=b&origin=searchbasic&rr=&sl=52&s=TITLE-ABS-KEY\(information%20security%20machine%20learning\)&searchterm1=information%20security%20machine%20learning&searchTerms=&connectors=](https://www2.scopus.com/results/results.uri?src=s&st1=&st2=&sot=b&sdt=b&origin=searchbasic&rr=&sl=52&s=TITLE-ABS-KEY(information%20security%20machine%20learning)&searchterm1=information%20security%20machine%20learning&searchTerms=&connectors=) Lekérdezés időpontja: 2021. 01. 01.
163. SCOPUS (2021c): [https://www.scopus.com/results/results.uri?src=s&st1=&st2=&sot=b&sdt=b&origin=searchbasic&rr=&sl=28&s=TITLE-ABS-KEY\(deep%20learning\)&searchterm1=deep%20learning&searchTerms=&connectors=](https://www.scopus.com/results/results.uri?src=s&st1=&st2=&sot=b&sdt=b&origin=searchbasic&rr=&sl=28&s=TITLE-ABS-KEY(deep%20learning)&searchterm1=deep%20learning&searchTerms=&connectors=) Lekérdezés időpontja: 2021. 01. 01.
164. SHABBIR, J. – ANWER, T. (2015): Artificial Intelligence and its Role in Near Future. In: *Journal of Latex Class Files*, 4 (8) 1-11. p.
165. SHANNON, C. E. (1948): A Mathematical Theory of Communication. In: *The Bell System Technical Journal*, (27) 379-423. p.
166. SHARMA, V. - RAVINDER, K. – CHENG, W. – ATIQUZZAMAN, M. – SRINIVASAN, K. – ZOMAYA, A. Y. (2018): NHAD: Neuro-Fuzzy Based Horizontal Anomaly Detection In Online Social Networks. In: *IEEE Transactions on Knowledge and Data Engineering*, 30 (11): 2171-2184. p.
167. SHEKHAR, S. – AKOGLU, L. (2019): Incorporating Privileged Information to Unsupervised Anomaly Detection. In: JOINT EUROPEAN CONFERENCE ON MACHINE LEARNING

- AND KNOWLEDGE DISCOVERY IN DATABASES (2018)(Ghent). Machine learning and knowledge discovery in databases: proceedings part I. Ghent, Springer. p. 1-16.
168. SHORTEN, C. – KHOSHGOFTAAR, T. M. (2019): A survey on Image Data Augmentation for Deep Learning. In: *Journal of Big Data*, 6 (1) 1-48. p.
 169. SMAHA, S. E. (1988): Haystack: An intrusion detection system. In: AEROSPACE COMPUTER SECURITY APPLICATIONS CONFERENCE (4.)(1988)(Orlando). IEEE Fourth Aerospace Computer Security Applications Conference: proceedings. Orlando, p. 37-44.
 170. STEGMAN E. – GUEVARA, J. – MICHELOGIKANNAKIS, N. – FUTELA, S. – SHARMA, S. – KAUSHAL, S. (2019): IT Key Metrics Data 2020: Industry Measures — Executive Summary. <https://www.gartner.com/document/3975995?ref=gfeed> Keresőprogram: Google. Kulcsszavak: IT Key metrics. Lekérdezés időpontja: 2020. 08. 02.
 171. SZALAY, ZS. (2009): Menedzsment információs rendszerek gazdasági elemzése. Doktori disszertáció. Gödöllő: Szent István Egyetem, Gazdálkodás- és Szervezéstudományi Doktori Iskola. 155 p.
 172. SZELÉNYI, L. (2001): Többváltozós ökonometriai módszerek. Gödöllő: Szent István Egyetem Kiadó. 103 p.
 173. SZEPESNÉ, S. M. (2010): Rendszertervezés 1. Az információrendszer fogalma, feladata, fejlesztése. Székesfehérvár: TÁMOP – 4.1.2-08/I/A-2009-0027, Nyugat-magyarországi Egyetem: Geoinformatikai Kar. 15 p.
 174. SZŰCS, I. (Szerk.) (2008): A tudományos megismerés rendszertana. Budapest: Szent István Egyetem Kiadó. 272 p.
 175. TAN, P. – STEINBACH, M. – KUMAR, V. et al. (2018): Introduction to Data Mining. Second Edition. USA: Pearson. 864 p.
 176. UGRÓSDY, GY. (2018): Gazdaságstatisztika. Gödöllő: Szent István Egyetemi Kiadó. 113 p.
 177. VARGA, J. – CSEH, B. (2019): A negyedik ipari forradalom egyes adózási és munkaerőpiaci hatásai. In: *Controller Info*, 7 (1) 11-14. p.
 178. VASVÁRI, GY. (2008): Vállalati (szervezeti) kockázatmenedzsment. Budapest: Információs Társadalomért Alapítvány. 183 p.
 179. VAUGHAN, J. – SUDJANTO, A. – BRAHIMI, E. – CHEN, J. – NAIR, V. N. (2018): Explainable Neural Networks based on Additive Index Models. <https://arxiv.org/pdf/1806.01933.pdf> Keresőprogram: Google. Kulcsszavak: explainable neural network. Lekérdezés időpontja: 2020. 09. 24.
 180. VINOGRADOV, SZ. (2020): A nemzeti versenyképesség puha tényezői, a társadalmi versenyképesség. 109-138. p. In: CSATH M. (Szerk.): *Versenyképesség: új elméleti és módszertani megközelítések*. Budapest: Dialóg Campus Kiadó, 215 p.
 181. WADE, C. (2020): Hands-On Gradient Boosting with XGBoost and scikit-learn: Perform accessible machine learning and extreme gradient boosting with Python. Birmingham: Packt Publishing. 310 p.
 182. WARD, J. (1998): Az információ-rendszerek szervezési elvei. Budapest: Co-Nex Könyvkiadó. 243 p.
 183. WARRENDER C. – FORREST, S. – PEARLMUTTER, B. (1999): Detecting intrusions using system calls: Alternative data models. In: SYMPOSIUM ON SECURITY AND PRIVACY (1999)(Oakland). IEEE Symposium on Security and Privacy. Oakland, p. 133-145.
 184. Wikiszótár (2020): <https://wikiszotar.hu/ertelmezo-szotar/Gener%C3%A1l> Keresőprogram: Google. Kulcsszavak: wikiszótár generál. Lekérdezés időpontja: 2020. 08. 17.

185. WINSTON, P. H. (1992): Artificial Intelligence. Third edition. Massachusetts: Addison-Wesley. 737 p. Idézve: RUSSEL, S. – NORVIG, P. (2005): Mesterséges Intelligencia modern megközelítésben. Budapest: Panem Kiadó. 1206 p.
186. WOLPERT, D. H. (2020): What is important about the No Free Lunch theorems? <https://arxiv.org/pdf/2007.10928.pdf> Keresőprogram: Google. Kulcsszavak: no free lunch theorem. Lekérdezés időpontja: 2020. 12. 22.
187. WORTH, S. – GROSS, L. (1977): Szimbolikus stratégiák. 37-50. p. In: BUSIGNIES, H. – STENT, G. S. – KNAPP, M. L. (Szerk.): *Kommunikáció I. Válogatott tanulmányok*. Budapest: Közgazdasági és Jogi Kiadó, 317 p.
188. YAMANISHI, K. – TAKEUCHI, J. I. (2001): Discovering outlier filtering rules from unlabeled data: combining a supervised learner with an unsupervised learner. In: INTERNATIONAL CONFERENCE ON KNOWLEDGE DISCOVERY AND DATA MINING (7.)(2001)(San Francisco). Proceedings of the Seventh ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. San Francisco, p. 389-394.
189. YANG, F. – ZHANG, W. – TAO, L. – MA, J. (2020): Transfer Learning Strategies for Deep Learning-based PHM Algorithms. In: *Applied Sciences*, 10 (7) 1-19. p.
190. YE, N. – LI, X. – CHEN, Q. – EMRAN, S. M. – XU, M. (2001): Probabilistic techniques for intrusion detection based on computer audit data. In: *IEEE Transactions on Systems, Man, and Cybernetics – Part A: Systems and Humans*, 31 (4) 266-274. p.
191. YE, N. – EMRAN, S. M. – CHEN, Q. – VILBERT, S. (2002): Multivariate statistical analysis of audit trails for host-based intrusion detection. In: *IEEE Transaction on Computers*, 810-820. p.
192. YEMM, G. (2012): Leading Your Team: How to Set Goals, Measure Performance and Reward Talent. London: Pearson. 171 p.
193. YOUNG, L. (2020): Risk IT Framework. 2nd Edition. USA: ISACA. 46 p.
194. YU, B. – KUMBIER, K. (2018): Artificial Intelligence and Statistics. In: *Frontiers of Information Technology & Electronic Engineering*, 19 (1) 6-9. p.
195. YU, E. – PAREKH, P. (2016): A Bayesian Ensemble for Unsupervised Anomaly Detection. <https://arxiv.org/pdf/1610.07677.pdf>. Keresőprogram: Google. Kulcsszavak: event and anomaly detection. Lekérdezés időpontja: 2020. 02. 22.
196. ZOLTAYNÉ, P. R. (2002): Döntéelmélet. Budapest: Alinea Kiadó. 596 p.
197. ZHANG, J. – ZULKERNIE, M. (2006a): A hybrid network intrusion detection technique using random forests. In: INTERNATIONAL CONFERENCE ON AVAILABILITY (1.)(2006)(Vienna). Proceedings of the First International Conference on Availability, Reliability and Security. Vienna, 262-269. p.
198. ZHANG, J. – ZULKERNIE, M. (2006b): Anomaly based network intrusion detection with unsupervised outlier detection. In: INTERNATIONAL CONFERENCE ON COMMUNICATIONS (2006)(Istanbul). IEEE International Conference on Communications: proceedings. Istanbul, 2388-2393 p.
199. ZHENG, A. – CASARI, A. (2018): Feature Engineering for Machine Learning. USA: O'Reilly. 200 p.

Jogszabályi hivatkozások

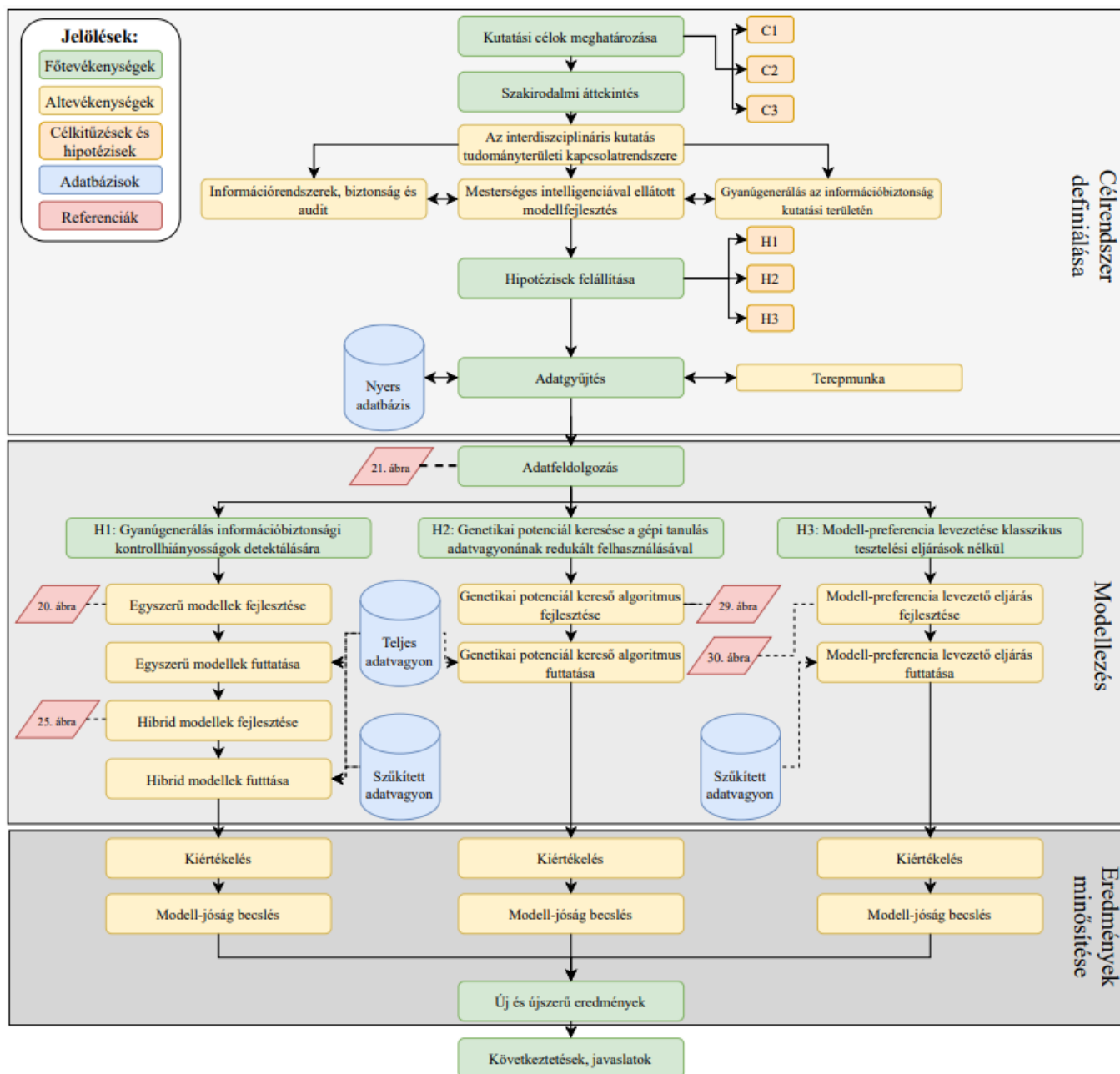
- 200.** 19/2017. (VII. 19.) MNB rendelet a pénzmosás és a terrorizmus finanszírozása megelőzéséről és megakadályozásáról szóló törvény végrehajtásának az MNB által felügyelt szolgáltatókra vonatkozó, valamint az Európai Unió és az ENSZ Biztonsági Tanácsa által elrendelt pénzügyi és vagyoni korlátozó intézkedések végrehajtásáról szóló törvény szerinti szűrőrendszer kidolgozásának és működtetése minimumkövetelményeinek részletes szabályairól.
- 201.** 42/2015. (III. 12.) Korm. rendelet a pénzügyi intézmények, a biztosítók és a viszontbiztosítók, továbbá a befektetési vállalkozások és az árutőzsdei szolgáltatók informatikai rendszerének védelméről.
- 202.** 45/2018. (XII. 17.) MNB rendelet a pénzmosás és a terrorizmus finanszírozása megelőzéséről és megakadályozásáról szóló törvény végrehajtásának az MNB által felügyelt szolgáltatókra vonatkozó, valamint az Európai Unió és az ENSZ Biztonsági Tanácsa által elrendelt pénzügyi és vagyoni korlátozó intézkedések végrehajtásáról szóló törvény szerinti szűrőrendszer kidolgozásának és működtetése minimumkövetelményeinek részletes szabályairól.
- 203.** 26/2020. (VIII. 25.) MNB rendelet a pénzmosás és a terrorizmus finanszírozása megelőzéséről és megakadályozásáról szóló törvény végrehajtásának az MNB által felügyelt szolgáltatókra vonatkozó, valamint az Európai Unió és az ENSZ Biztonsági Tanácsa által elrendelt pénzügyi és vagyoni korlátozó intézkedések végrehajtásáról szóló törvény szerinti szűrőrendszer kidolgozásának és működtetése minimumkövetelményeinek részletes szabályairól.
- 204.** 2017. évi LIII. törvény a pénzmosás és a terrorizmus finanszírozása megelőzéséről és megakadályozásáról.
- 205.** A Magyar Nemzeti Bank 2/2017. (I.12.) számú ajánlása a közösségi és publikus felhőszolgáltatások igénybevételéről.
- 206.** A Magyar Nemzeti Bank 4/2019. (IV.1.) számú ajánlása a közösségi és publikus felhőszolgáltatások igénybevételéről.
- 207.** AZ EURÓPAI PARLAMENT ÉS A TANÁCS (EU) 2016/679 RENDELETE (2016. április 27.) a természetes személyeknek a személyes adatok kezelése tekintetében történő védelméről és az ilyen adatok szabad áramlásáról, valamint a 95/46/EK rendelet hatályon kívül helyezéséről (általános adatvédelmi rendelet - GDPR).
- 208.** Vezetői körlevél az elektronikus úton megkötött írásbeli szerződésekről, megtett írásbeli jognyilatkozatokról.

Szabványok

- 209.** ISO/IEC 27000 (2014): Information technology – Security techniques – Information security management systems – Overview and vocabulary. International Standard. 31 p.
- 210.** ISO/IEC 27001 (2013): Information technology – Security techniques – Information security management systems – Requirements. International Standard. 23 p.
- 211.** National Institute of Standards and Technology (2013): Security and Privacy Controls for Federal Information Systems and Organizations. USA: NIST Special Publication 800-53. 462 p.

10. MELLÉKLETEK

1. sz. melléklet: A kutatás folyamatábrája



2. sz. melléklet: Az ISO/IEC 27001:2013 A melléklet kontrollterületeinek rövid bemutatása

A5. Az információbiztonság vezetői irányítása: Menedzsment direktívák, szabályzatok és utasítások elkészítésére és fenntartására vonatkozó követelmények.

A6. Az információbiztonság szervezete: A terület az információbiztonság szervezeti keretrendszerét, határait és felelősségköreit határozza meg taktikai és operatív szinten.

A7. Humán-erőforrás biztonsága: A szervezet humán-erőforrás biztonságának alapelveit, a szervezet munkatársaira és partnereire vonatkozó elvárásokat fekteti le, beleértve a munkavállalók átvilágítását, információbiztonságra vonatkozó kötelelességeit, folyamatos oktatását és fegyelmi eljárásokat információbiztonsági kihágások és a szabályzatokban lefektetett követelmények megsértése esetén.

A8. Vagyon-menedzsment: A szervezet vagyontárgyaira és eszközkezelésére vonatkozó követelmények összességét fejt ki, mely magában foglalja a vagyonelemek teljeskörűségének dokumentálását, naprakészen tartását és kapcsolódó felelősségi köröket. Részletezi az adatok biztonsági osztályba sorolásának követelményrendszerét, valamint az adathordozók biztonságos használatát, leselejtezését és szállítását.

A9. Hozzáférés szabályozás: Az információs vagyonhoz történő hozzáférések menedzsmentjével és szabályozásával foglalkozó kontrollterület. Taglalja az informatikai erőforrásokhoz történő jogosultságok megadásának, módosításának, megvonásának és felülvizsgálatának elvárásait kitérve a privilegizált, generikus és technikai felhasználókra egyaránt. Továbbá, követelményeket definiál a biztonságos bejelentkezés kikényszerítéséhez és jelszókezeléshez.

A10. Titkosítás: A kriptográfiai folyamatok alappilléreit definiálja, hangsúlyozva a titkosítási kulcsok és tanúsítványok menedzsmentjét.

A11. Fizikai és környezeti biztonság: A szervezet biztonsági zónáira (pl. irodaház, szerverterem, stb.) és területeire vonatkozó utasításokat összegez, úgy, mint az épületek környezeti kontrollokkal (pl. tűzoltóberendezés, páratartalom mérés stb.) történő ellátása, a kábelezés biztonsági kialakítása és a védett helyekhez történő hozzáférések ellenőrzése, valamint részletezi a fizikai eszközök megfelelő használatát.

A12. A működtetés biztonsága: Az informatikai üzemeltetéssel kapcsolatos független biztonsági kontrollok követelményrendszerét azonosítja, melybe beleértendő a változáskezelés szabályozása, a kapacitások monitorozása, a vírusok és más rosszindulatú programok elleni védelem felállítása, az adatok mentése és visszaállítása, a naplózási folyamatok kialakítása és független felülvizsgálata, továbbá a technikai sérülékenységek azonosítása és elhárítása.

A13. A kommunikáció biztonsága: A hálózatokon történő kommunikáció és adattranszfer követelményeit fekteti le, részletezi a hálózatok biztonságos szeparációját és a kommunikációhoz (pl. elektronikus levelezéshez) szükséges védelmi intézkedéseket.

A14. Rendszer beszerzés, fejlesztés és karbantartás: Az információbiztonság beépítését részletezi az informatikai rendszerek teljes életciklusára kivetítve. A beszerzésre, fejlesztésre és karbantartásra vonatkozó kontrollkövetelmények kitérnek a bizalmas adatok kezelésére, a biztonságos fejlesztői környezet üzemeltetésére és tesztelési adatok körültekintő felhasználására.

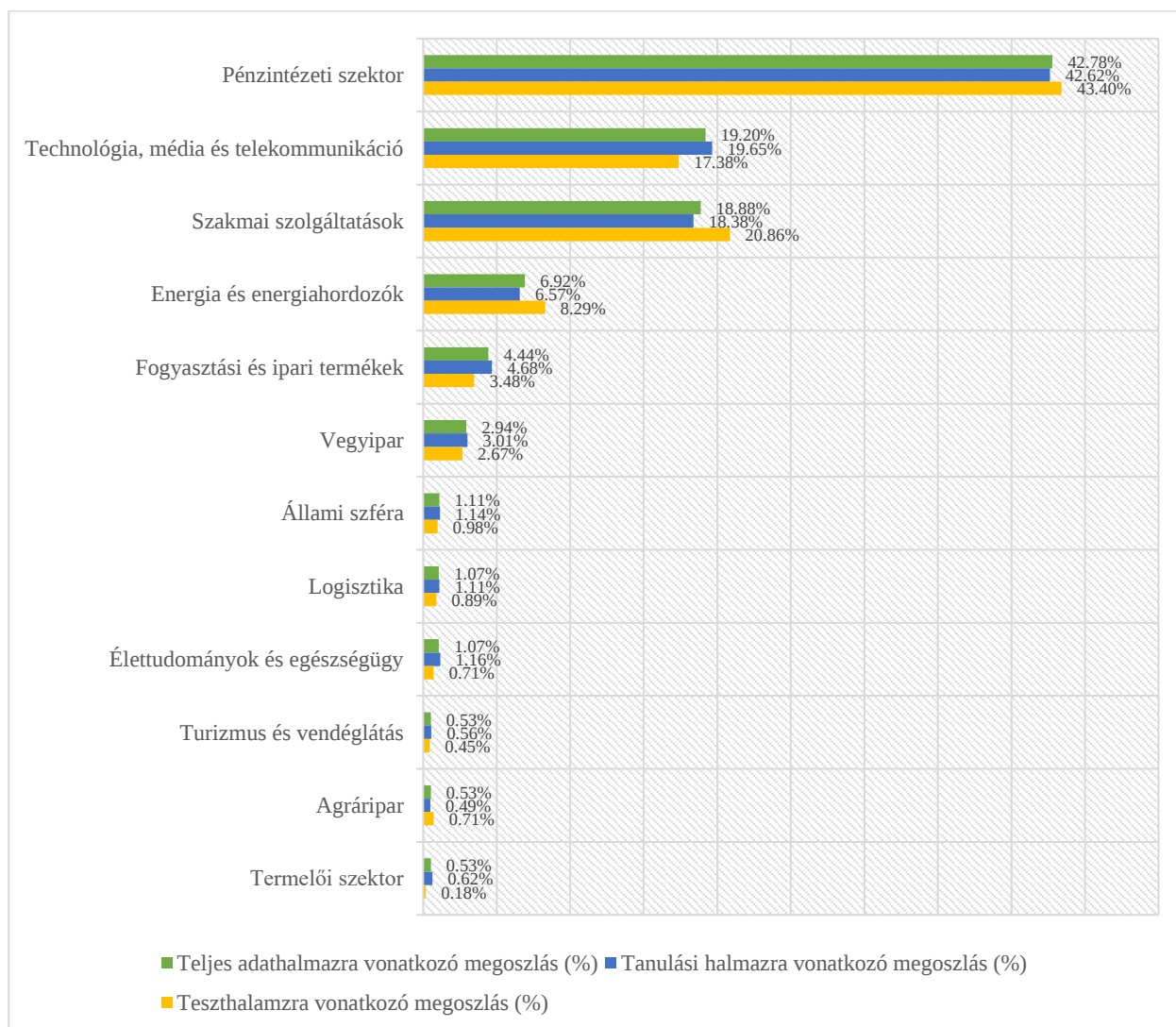
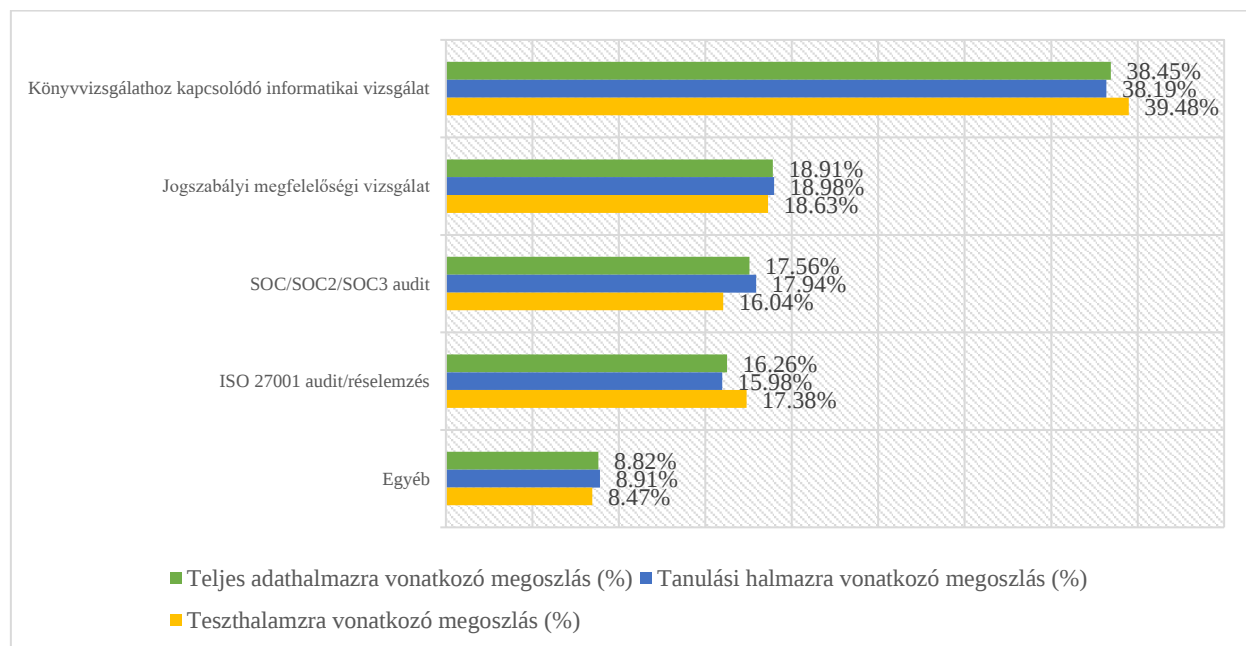
A15. Szállítói kapcsolatok: A szállítók kezelésével kapcsolatos követelményrendszert definiálja, melyben hangsúlyt fektet a beszállítók biztonsági követelményeknek történő megfelelésére (pl. külső hozzáférések szabályozása az informatikai erőforrásokhoz), a szállítók auditálására, szolgáltatásmenedzsmentjére és szerződéses követelmények betartatására.

A16. Információbiztonsági incidensek kezelése: A biztonsági incidensek menedzsmentjét tárgyalja, a munkavállalók felelősségét az incidensek jelentésére, az incidensek megszüntetésének módjait, gyökereinek feltérképezését és alátámasztó evidenciák begyűjtését/megőrzését, valamint az incidensek dokumentálását és azokból történő tudás felhasználását a következő hasonló incidens megakadályozása érdekében.

A17. A működésfolytonosság információbiztonsági aspektusai: Az üzletmenetfolytonosság fenntartását célzó kontrollkövetelmények összessége, mely magában foglalja az üzleti hatáselemzések kivitelezését, a folytonosságra vonatkozó követelményrendszer felállítását az üzleti hatáselemzés tükrében, valamint az üzletmenetfolytonosságot szavatoló kontrollok folyamatos tesztelését és ellenőrzését.

A18. Megfelelőség: A külső (pl. releváns területi és regionális jogszabályok és rendeletek, szerződéses kötelezettségek, stb.) és belső (szabályzatok és irányelvek) követelményeknek történő megfelelés biztosítására vonatkozó követelményeket határoz meg.

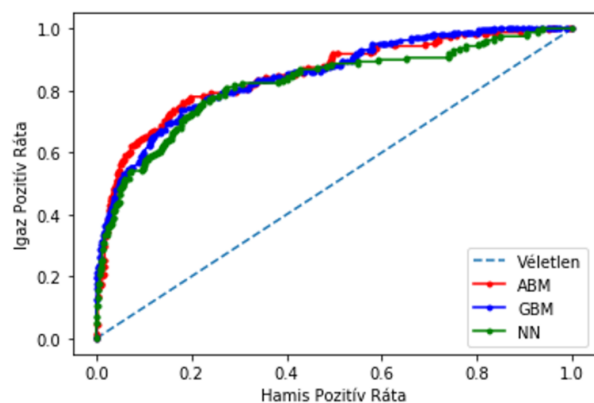
3. sz. melléklet: A feldolgozott adatvagyon megoszlásai a teljes, tanuló-, valamint teszhalmazon audittípusonként és iparáganként.



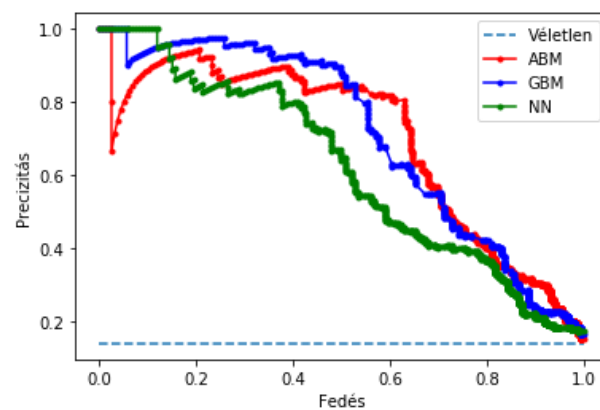
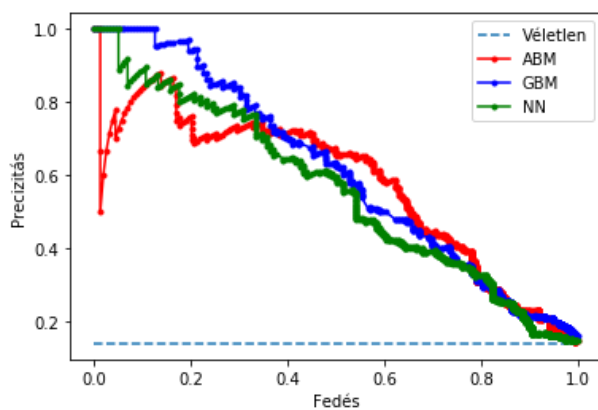
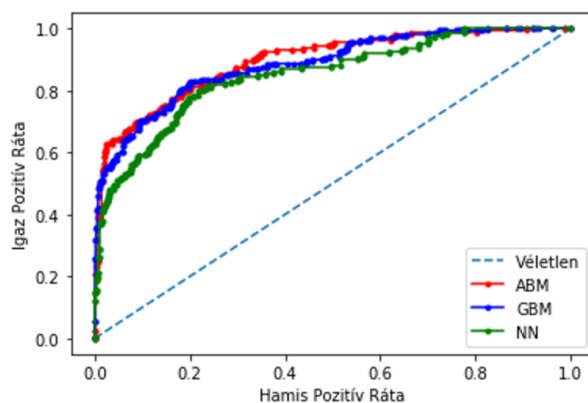
Forrás: Saját szerkesztés

4. sz. melléklet: Az egyszerű és hibrid modellek ROC és PR-görbéi

Egyszerű modellek

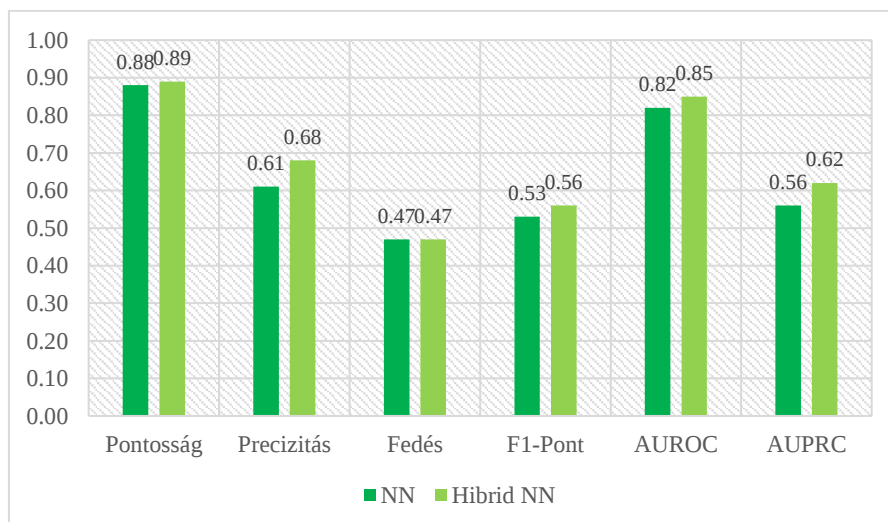
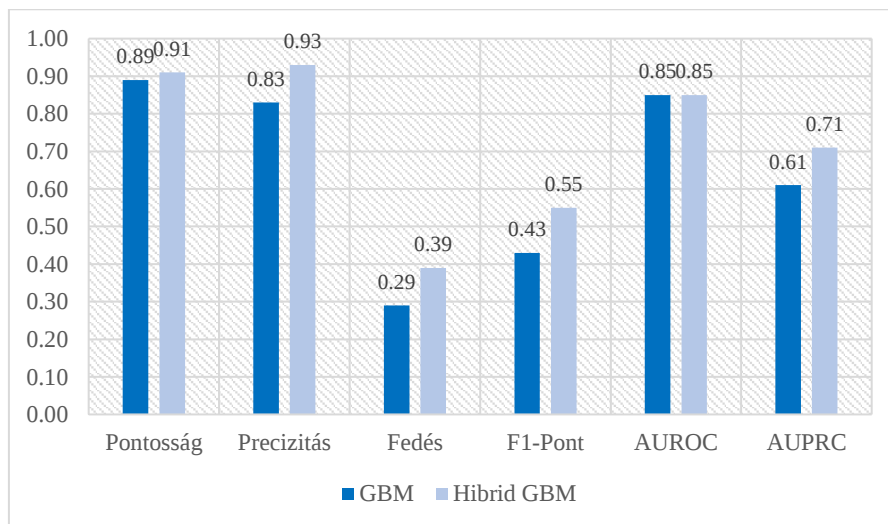
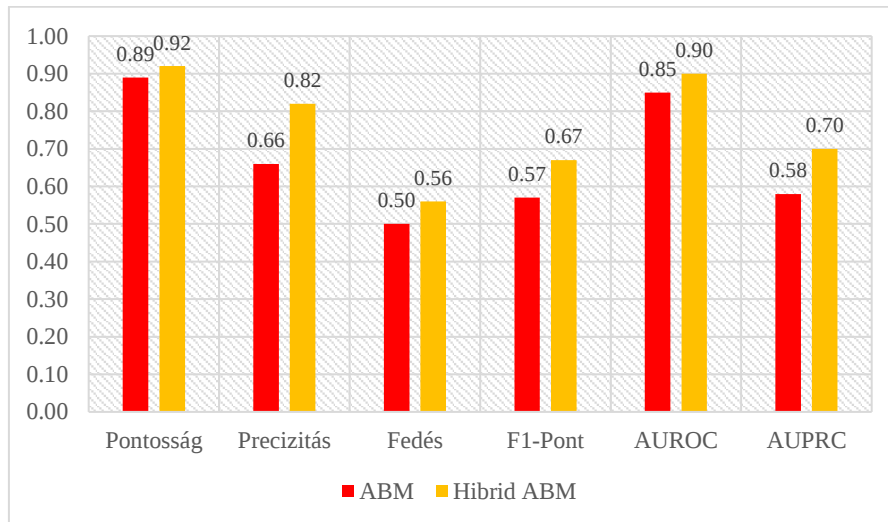


Hibrid modellek



Forrás: Saját szerkesztés

5. sz. melléklet: Egyszerű és hibrid modellek jószágmetrikái grafikusan oszlopdiagrammon



Forrás: Saját szerkesztés

6. sz. melléklet: Egyszerű és hibrid felügyelt tanuló modellek performancia metrikái iparági megoszlásban

<i>Iparágak</i>	Egyszerű ABM		Egyszerű GBM		Egyszerű NN	
	Pontosság	F1-Pont	Pontosság	F1-Pont	Pontosság	F1-Pont
<i>Agrárpar</i>	0.88	n/a	1.00	n/a	1.00	n/a
<i>Állami szféra</i>	1.00	n/a	1.00	n/a	1.00	n/a
<i>Élettudományok és egészségügy</i>	0.75	n/a	0.88	n/a	0.88	n/a
<i>Energia és energiahordozók</i>	0.89	0.74	0.86	0.55	0.88	0.69
<i>Fogyasztási és ipari termékek</i>	0.74	0.29	0.79	0.43	0.79	0.33
<i>Logisztika</i>	1.00	1.00	0.90	n/a	0.90	n/a
<i>Pénzüntézet i szektor</i>	0.90	0.52	0.90	0.42	0.87	0.44
<i>Szakmai szolgáltatások</i>	0.88	0.58	0.87	0.38	0.87	0.54
<i>Technológia, média és telekommunikáció</i>	0.92	0.50	0.92	0.38	0.93	0.63
<i>Termelői szektor</i>	1.00	n/a	1.00	n/a	1.00	n/a
<i>Turizmus és vendéglátás</i>	1.00	1.00	0.80	n/a	0.80	n/a
<i>Vegyipar</i>	0.90	0.82	0.83	0.62	0.90	0.80
<i>Teljes adatvagyón</i>	0.89	0.57	0.89	0.43	0.88	0.53

<i>Iparágak</i>	Hibrid ABM		Hibrid GBM		Hibrid NN	
	Pontosság	F1-Pont	Pontosság	F1-Pont	Pontosság	F1-Pont
<i>Agrárpar</i>	1.00	n/a	1.00	n/a	1.00	n/a
<i>Állami szféra</i>	1.00	n/a	1.00	n/a	1.00	n/a
<i>Élettudományok és egészségügy</i>	1.00	n/a	1.00	n/a	0.75	n/a
<i>Energia és energiahordozók</i>	0.96	0.89	0.90	0.71	0.87	0.71
<i>Fogyasztási és ipari termékek</i>	0.82	0.46	0.79	0.43	0.79	0.43
<i>Logisztika</i>	1.00	1.00	0.90	n/a	0.90	n/a
<i>Pénzüntézet i szektor</i>	0.93	0.67	0.92	0.54	0.90	0.55
<i>Szakmai szolgáltatások</i>	0.88	0.58	0.89	0.52	0.87	0.47
<i>Technológia, média és telekommunikáció</i>	0.93	0.53	0.92	0.38	0.92	0.50
<i>Termelői szektor</i>	1.00	n/a	1.00	n/a	1.00	n/a
<i>Turizmus és vendéglátás</i>	1.00	1.00	0.80	n/a	0.80	n/a
<i>Vegyipar</i>	0.93	0.86	0.93	0.86	0.97	0.93
<i>Teljes adatvagyón</i>	0.92	0.67	0.91	0.55	0.89	0.56

Forrás: Saját szerkesztés

7. sz. melléklet: Felügyelt egyszerű és hibrid módszerek performancia metrikái szűkített és a teljes adatvagyonon

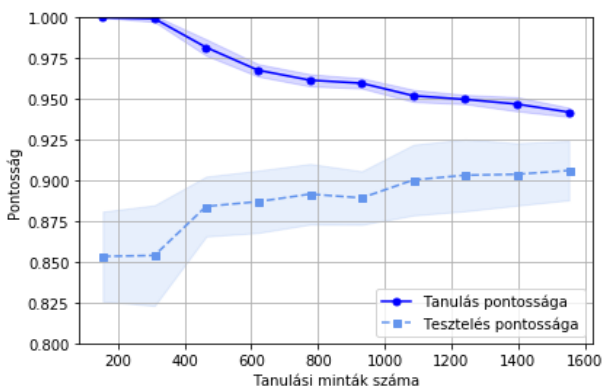
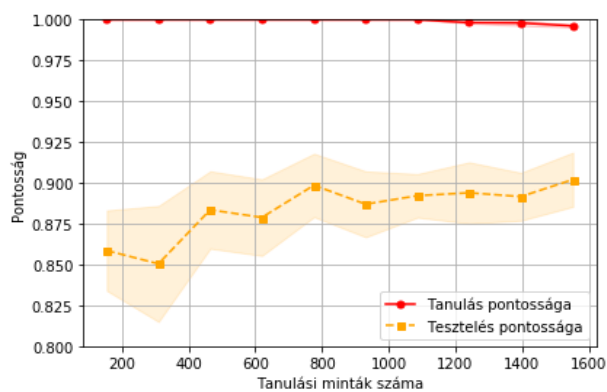
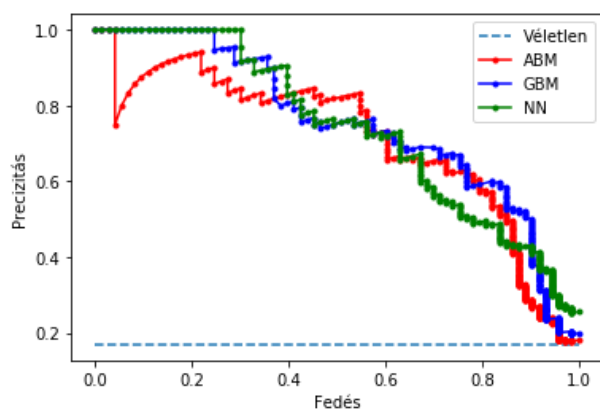
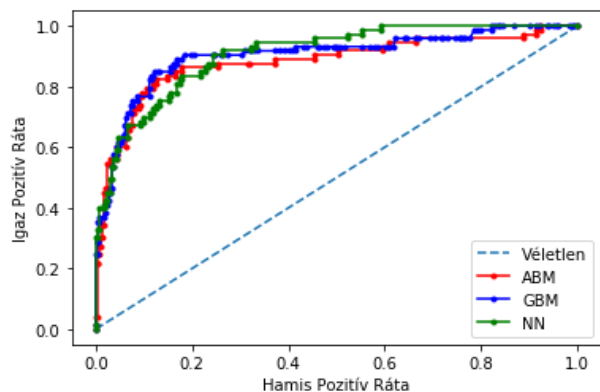
Performancia mutatók	Szűkített adatvagyon (pénzüintézet szektor)			Teljes adatvagyon (csak pénzüintézet szektor eredményei)		
	<i>Egyszerű ABM</i>	<i>Egyszerű GBM</i>	<i>Egyszerű NN</i>	<i>Egyszerű ABM</i>	<i>Egyszerű GBM</i>	<i>Egyszerű NN</i>
<i>Igaz pozitívák száma (db)</i>	29	21	23	28	17	24
<i>Igaz negatívák száma (db)</i>	393	409	393	408	423	401
<i>Hamis pozitívák száma (db)</i>	22	6	22	19	4	26
<i>Hamis negatívák száma (db)</i>	36	44	42	32	43	36
<i>Pontosság</i>	0.88	0.90	0.87	0.90	0.90	0.87
<i>Precizitás</i>	0.57	0.78	0.51	0.60	0.81	0.48
<i>Fedés</i>	0.45	0.32	0.35	0.47	0.28	0.40
<i>F1-Pont</i>	0.50	0.46	0.42	0.52	0.42	0.44
<i>Variancia</i>	0.10	0.05	0.09	n/a	n/a	n/a
<i>AUROC</i>	0.78	0.79	0.75	n/a	n/a	n/a
<i>AUPRC</i>	0.49	0.52	0.38	n/a	n/a	n/a

Performancia mutatók	Szűkített adatvagyon (pénzüintézet szektor)			Teljes adatvagyon (csak pénzüintézet szektor eredményei)		
	<i>Hibrid ABM</i>	<i>Hibrid GBM</i>	<i>Hibrid NN</i>	<i>Hibrid ABM</i>	<i>Hibrid GBM</i>	<i>Hibrid NN</i>
<i>Igaz pozitívák száma (db)</i>	29	25	21	35	23	29
<i>Igaz negatívák száma (db)</i>	404	412	397	417	425	410
<i>Hamis pozitívák száma (db)</i>	11	3	18	10	2	17
<i>Hamis negatívák száma (db)</i>	36	40	44	25	37	31
<i>Pontosság</i>	0.90	0.91	0.87	0.93	0.92	0.90
<i>Precizitás</i>	0.73	0.89	0.54	0.78	0.92	0.63
<i>Fedés</i>	0.45	0.38	0.32	0.58	0.38	0.48
<i>F1-Pont</i>	0.55	0.54	0.40	0.67	0.54	0.55
<i>Variancia</i>	0.08	0.04	0.09	n/a	n/a	n/a
<i>AUROC</i>	0.80	0.82	0.76	n/a	n/a	n/a
<i>AUPRC</i>	0.55	0.63	0.40	n/a	n/a	n/a

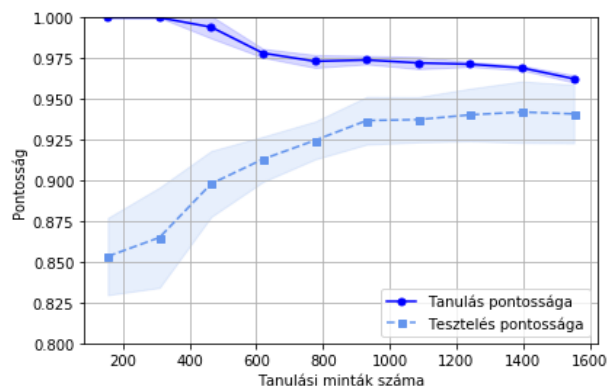
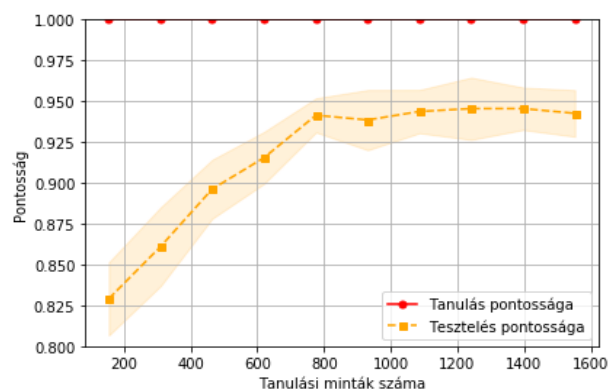
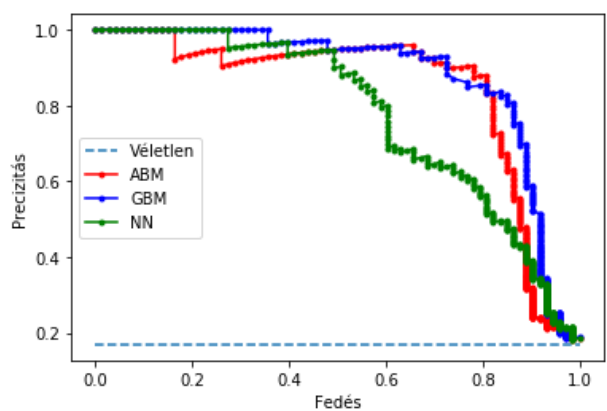
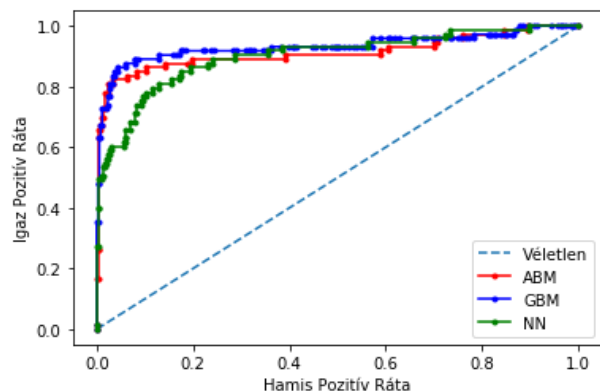
Forrás: Saját szerkesztés

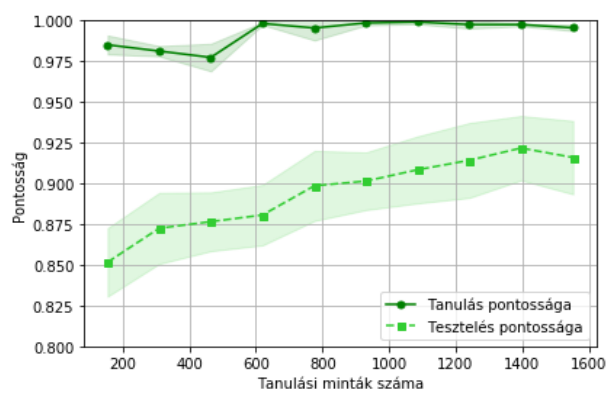
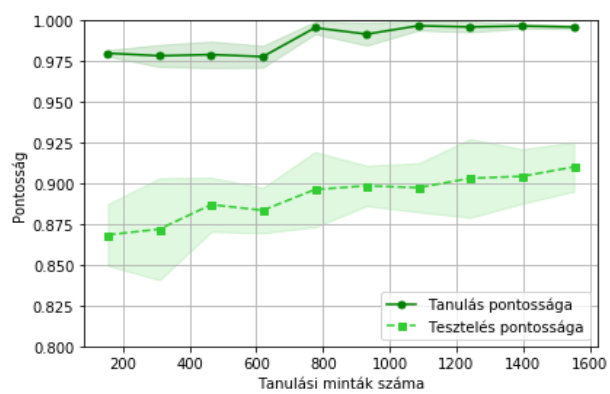
8. sz. melléklet: Egyszerű és hibrid felügyelt módszerek ROC és PR-görbéi, valamint tanulási görbéi szűkített adatvagyonon (könyvvizsgálathoz kapcsolódó informatikai vizsgálat)

Egyszerű modellek



Hibrid modellek

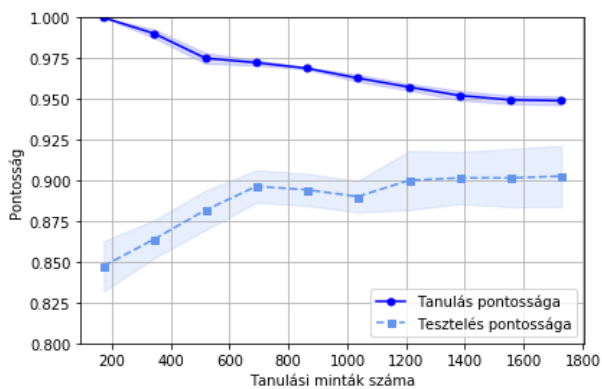
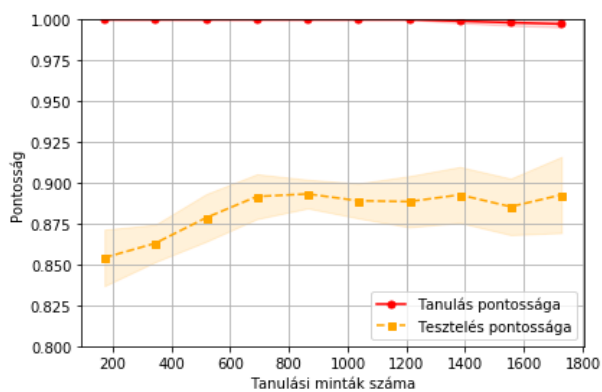
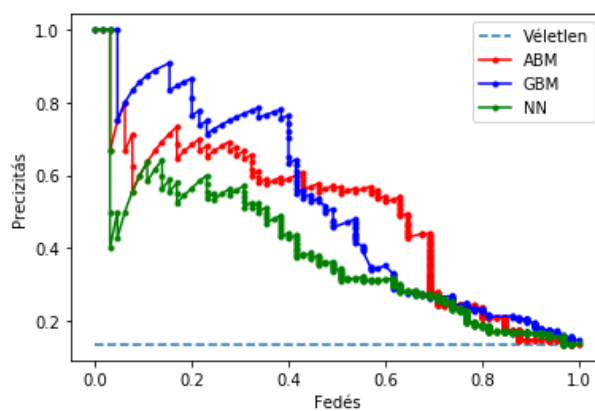
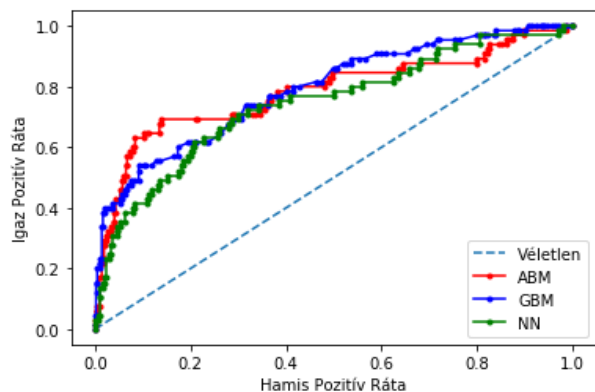




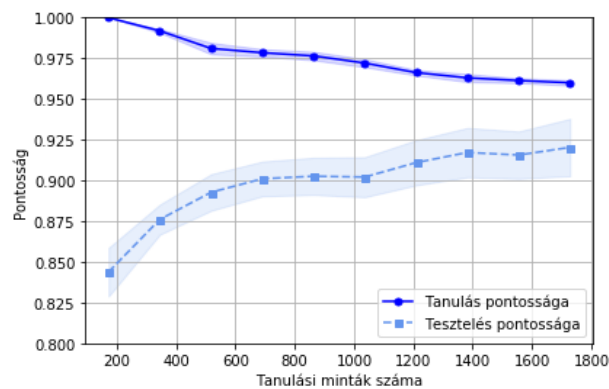
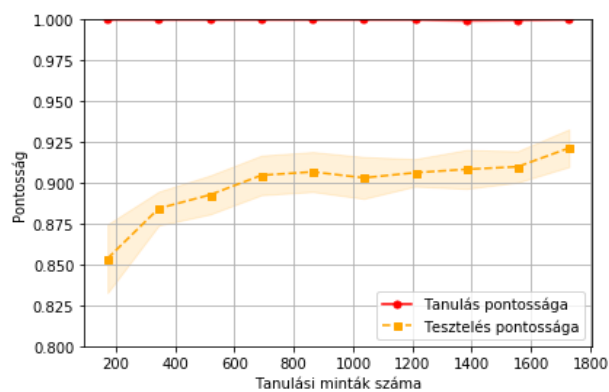
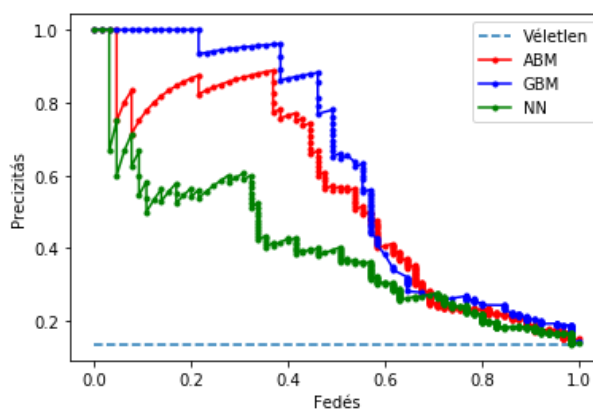
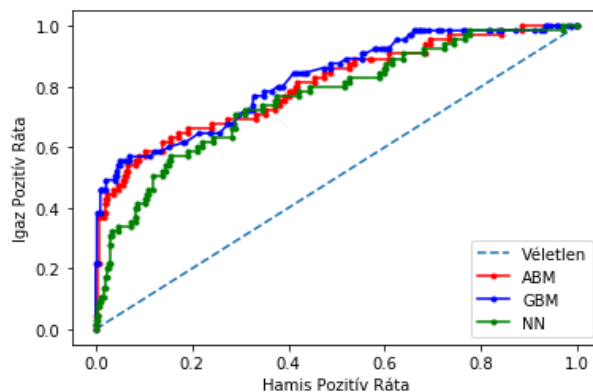
Forrás: Saját szerkesztés

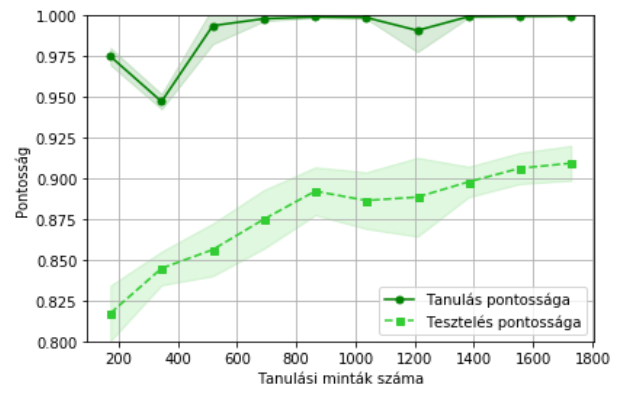
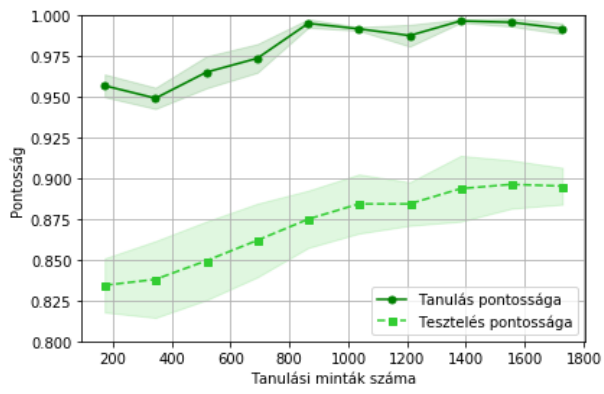
9. sz. melléklet: Egyszerű és hibrid felügyelt módszerek ROC-görbéi és tanulási görbéi szűkített adatvagyonon (pénzügyeti szektor)

Egyszerű modellek



Hibrid modellek





Forrás: Saját szerkesztés

10. sz. melléklet: Felügyelt modellek kategóriáinak értékelése varianciaelemzéssel

Varianciaelemzésben felhasznált és kiértékelt változók	Modellkomplexitás szerinti kategóriák	
<i>Idealitás mutató</i>	Hasonlóságelemzés	
<i>Kategória</i>	Egyszerű modellek	Hibrid modellek
<i>Felhasznált eset</i>	9	9
<i>Átlag</i>	986.81	1103.22
<i>Szórás</i>	28.90	34.38
<i>Levene-teszt szignifikanciája</i>	0.67	
<i>F-próba</i>	3.11	
<i>Szignifikancia</i>	0.10	

Varianciaelemzésben felhasznált és kiértékelt változók	Alkalmazott módszerek szerinti kategóriák		
<i>Idealitás mutató</i>	Hasonlóságelemzés		
<i>Kategória</i>	ABM	GBM	NN
<i>Felhasznált eset</i>	6	6	6
<i>Átlag</i>	1004.78	1011.55	983.72
<i>Szórás</i>	32.76	30.26	36.89
<i>Levene-teszt szignifikanciája</i>	0.75		
<i>F-próba</i>	1.13		
<i>Szignifikancia</i>	0.35		

Varianciaelemzésben felhasznált és kiértékelt változók	Felhasznált adatvagyon szerinti kategóriák		
<i>Idealitás mutató</i>	Hasonlóságelemzés		
<i>Kategória</i>	TA	KKIV	PSZ
<i>Felhasznált eset</i>	6	6	6
<i>Átlag</i>	993.73	1039.95	973.37
<i>Szórás</i>	26.64	19.21	23.98
<i>Levene-teszt szignifikanciája</i>	0.42		
<i>F-próba</i>	9.98		
<i>Szignifikancia</i>	0.00		
<i>Scheffé-próba szignifikanciája (TA – KKIV)</i>	0.04		
<i>Scheffé-próba szignifikanciája (TA – PSZ)</i>	0.35		
<i>Scheffé-próba szignifikanciája (PSZ – KKIV)</i>	0.00		

Forrás: Saját szerkesztés

11. sz. melléklet: A genetikai potenciál kereső és modell-preferencia levezetésre használt eljárások pszeudokódjai

Jelölések:

$\mathbf{A}$:= a tanulási halmaz tulajdonságait leíró mátrix

$\bar{\mathbf{A}}$:= Az $\mathbf{A}$ mátrix rangsorolással transzformált mátrixa (három dimenzió esetén az egyes irány-preferencia alternatívákat tartalmazó mátrixok vektora)

$\bar{\mathbf{A}}_{\text{counter}}$:= Az $\bar{\mathbf{A}}$ mátrix(ok) antri-diszkriminatív eljárással meghatározott adott konstans célváltozónál nagyobb értékkel rendelkező becsléseinek számát tartalmazó vektor

$\check{\mathbf{A}}$:= Az $\mathbf{A}$ (transzformált) mátrix modell-szintű aggregált mátrixa

$\tilde{\mathbf{A}}$:= Az $\check{\mathbf{A}}$ mátrix rangsorolással transzformált mátrixa

$\mathbf{I}$:= iránypreferenciákat tartalmazó vektor

$\mathbf{I}_{\text{agg}}$:= az aggregált mátrix iránypreferenciáit tartalmazó vektor

$\mathbf{M}$:= gépi tanuló modelleket tartalmazó vektor

$\mathbf{S}$:= termelési/anti-diszkriminatív függvény által szolgáltatott súlyokat tartalmazó mátrix

$\mathbf{X}$:= eredeti tanulóhalmaz

$\mathbf{X}_{\text{test}}$:= teszhalmaz

$\mathbf{Y}$:= célváltozó értékek vektora

$\tilde{\mathbf{Y}}$:= becsült célváltozó értékek vektora

Θ := konstans célváltozók vektora (pl. anti-diszkriminatív számításnál 1,000 értéket tartalmazó vektor)

abs(param₁) := abszolútértéket számoló eljárás.

aggregate(param₁, param₂, param₃(optional), param₄) := modell-szintre történő aggregáló eljárás, melynek első három bemenete 3 objektumszintű mátrix (a harmadik paraméter értelemszerűen opcionális), az utolsó paramétere a mátrixokra vonatkozó érvényes irány-preferenciákat tartalmazó vektor.

anti_discriminative_function(param₁, param₂) := anti-diszkriminatív függvényt illesztő eljárás, melynek két paramétere egy rangsorolt mátrix pl. $\bar{\mathbf{A}}$, valamint a célváltozó értékek Θ , kimenete a becsült célváltozók vektora $\tilde{\mathbf{Y}}$. (lásd 3.2.6. alfejezet)

append(param₁, param₂) := új objektum hozzáfűzése egy meglévő vektorhoz/mátrixhoz.

arg(param₁, param₂) := adott vektor/mátrix adott argumentumának pozíciójával visszatérő eljárás.

arg_max(param₁, param₂) := Adott vektor/mátrix adott argumentumához tartozó maximális érték.

check(param₁) := irány-preferenciák korrelációjának ellenőrzése, amennyiben közelít 0-hoz, igaz értékkel tér vissza.

compare(param₁, param₂, param₃, param₄) := az érvényes irány-preferencia nézeteket meghatározó eljárás, melynek első két paramétere két mátrix, különböző irány-preferenciaákkal (pl. $\bar{A}$), második két paramétere az antri-diszkriminatív eljárással meghatározott adott konstans célváltozónál nagyobb értékkel rendelkező becsléseinek számát tartalmazó vektorok ($\bar{A}_{\text{counter}}$). Visszatérési értéke az érvényes irány-preferencia nézeteket tartalmazó vektor.

correl(param₁, param₂) := Pearson-féle korreláció, melynek inputjai két vektor.

count(param₁) := adott vektor értékeit összeszámoló eljárás.

len(param₁) := adott paraméter hosszával visszatérő eljárás.

max(param₁) := adott vektor értékei közül a maximummal visszatérő eljárás.

norm(param₁, param₂) := normalizálást elvégző eljárás, melynek első paramétere egy mátrix, második az irány-preferenciákat tartalmazó vektor.

percentrank(param₁, param₂) := percentilis rangsorolást elvégző eljárás, melynek első paramétere egy mátrix, második az irány-preferenciákat tartalmazó vektor.

production_function(param₁, param₂) := termelési függvényt illesztő eljárás (lásd 3.2.6. alfejezet), mely első paramétere egy mátrix, második a mátrixhoz tartozó célváltozókból képzett vektor. Visszatérési értéke az **S** súlymátrix.

rank(param₁, param₂) := rangsorolást elvégző eljárás, melynek első paramétere egy mátrix, második az irány-preferenciákat tartalmazó vektor.

reduce(param₁, param₂, param₃) := tanulóhalmazt redukáló eljárás, melynek első paramétere a kiinduló tanulóhalmaz, második paramétere a tanulási halmaz tulajdonságait leíró mátrix adott iterációban kiválasztott legjobb attribútuma (aktuális és rákövető súlyszám különbsége), harmadik paramétere a kívánt súlyszámot elérni kívánt érték. A redukált tanulóhalmazzal tér vissza.

run(param₁) := gépi tanuló eljárás futtatása egy adott mátrixon (konstans konfigurációt feltételezve), melynek visszatérési értékei rendre: a tanulási halmaz tulajdonságait leíró új vektor, valamint a futtatáshoz tartozó célváltozó értéke.

sum(param₁) := adott vektor értékeit összegző eljárás.

Genetikai potenciál kereső eljárás (4.3. fejezet – 29. ábra)	
1:	procedure Gen-Pot-Search(X, X_test, A, Y)
2:	$X \leftarrow X_1 \times X_2$
3:	$X_test \leftarrow X_test_1 \times X_test_2$
4:	$A \leftarrow A_1 \times A_2$
5:	$Y \leftarrow Y_1$
6:	initialize $I \leftarrow I_1$
7:	initialize $\bar{A} \leftarrow \bar{A}_1 \times \bar{A}_2$
8:	initialize $S \leftarrow S_1 \times S_2$
9:	repeat
10:	$i \leftarrow i + 1, j \leftarrow j + 1, k \leftarrow k + 1$
11:	for i do until $\text{sum}(\max(S_1)) \leq \max(Y)$ or $\text{check}(I) = \text{true}$
12:	for j do until $j \leq \text{len}(A(j))$
13:	if $\text{correl}(A(j), Y) \geq 0$ then
14:	$I(j) \leftarrow 0$
15:	else $I(j) \leftarrow 1$
16:	end for
17:	$\bar{A} \leftarrow \text{rank}(A, I)$
18:	$S \leftarrow \text{production_function}(\bar{A}, Y)$
19:	$y \leftarrow \max(Y)$
20:	for k do until $k \leq \text{len}(S)$
21:	$l \leftarrow \text{arg}(S, S(k, \text{arg_max}(S, y)))$
22:	for l do until $l > 0$
23:	$\theta \leftarrow S(k, \text{arg_max}(S, y))$
24:	$\Delta \leftarrow \text{abs}(\theta - (\text{arg}(S, \theta) - l))$
25:	$l \leftarrow l - 1$
26:	if $\Delta > 0$ then
27:	break
28:	end for
29:	end for
30:	$D \leftarrow \text{append}(D, \Delta)$
31:	$\text{best_attribute} \leftarrow \text{arg_max}(S, \max(D))$
32:	$X \leftarrow \text{reduce}(X, \text{best_attribute}, \max(D))$
33:	$\Omega, y \leftarrow \text{run}(X_test)$
34:	$Y \leftarrow \text{append}(Y, y)$
35:	$A \leftarrow \text{append}(A, \Omega)$
36:	end for
37:	return X, y
38:	end procedure

Modell-preferencia levezető eljárás (4.4. fejezet – 30. ábra)	
1:	procedure Mod-Pre-Derivation(M, A)
2:	$M \leftarrow M_1$
3:	$A \leftarrow A_1 \times A_2$
4:	initialize $\bar{A} \leftarrow \bar{A}_1 \times \bar{A}_2 \times \bar{A}_3$
5:	initialize $\check{A} \leftarrow \check{A}_1 \times \check{A}_2$
6:	initialize $\bar{A_counter} \leftarrow \bar{A_counter}_1$
7:	initialize $\check{A} \leftarrow \check{A}_1 \times \check{A}_2$
8:	initialize counter
9:	initialize $I \leftarrow I_1$
10:	initialize $I_agg \leftarrow I_agg_1$
11:	repeat
12:	$i \leftarrow i + 1, j \leftarrow i + 1, k \leftarrow k + 1$
13:	for i do until $i \leq \text{len}(A_1)$
14:	for j do until $j \leq \text{len}(A_1)$
15:	if $\text{correl}(A(i), A(j)) \geq 0$ then
16:	$I(k) \leftarrow 0$
17:	else $I(k) \leftarrow 1$
18:	end for
19:	end for
20:	repeat
21:	$i \leftarrow i + 1$
22:	for i do until $i \leq \text{len}(A_1)$
23:	$\bar{A} \leftarrow \text{append}(\text{rank}(A, I(i)))$
24:	end for
25:	repeat
26:	$i \leftarrow i + 1, j \leftarrow i + 1, k \leftarrow k + 1$
27:	for i do until $i \leq \text{len}(\bar{A}_1)$
28:	for j do until $j \leq \text{len}(\bar{A}_1)$
29:	$\tilde{Y} \leftarrow \text{anti_discriminative_function}(\bar{A}(i), \Theta)$
30:	for k do until $k \leq \text{len}(\tilde{Y})$
31:	$\bar{A_counter} \leftarrow \text{append}(\bar{A_counter}, \text{count}(\tilde{Y}(k) > \Theta))$
32:	end for
33:	end for
34:	end for
35:	repeat
36:	$i \leftarrow i + 1, j \leftarrow i + 1$
37:	for i do until
38:	for j do until
39:	$\text{valid_views} \leftarrow \text{compare}(\bar{A}_3(i), \bar{A}_3(j), \bar{A_counter}(i), \bar{A_counter}(j))$
40:	end for
41:	end for
42:	$\check{A} \leftarrow \text{aggregate}(\bar{A}, \text{percentrank}(A, I), \text{norm}(A, I), \text{valid_views})$
43:	repeat
44:	$i \leftarrow i + 1, j \leftarrow i + 1, k \leftarrow k + 1$
45:	for i do until $i \leq \text{len}(\check{A})$
46:	for j do until $j \leq \text{len}(\check{A})$
47:	if $\text{correl}(\check{A}(i), \check{A}(j)) \geq 0$ then
48:	$I_agg(k) \leftarrow 0$
49:	else $I_agg(k) \leftarrow 1$

50:	end for
51:	end for
52:	$\tilde{A} \leftarrow \text{rank}(\tilde{A}, I_{\text{agg}})$
53:	$\tilde{Y} \leftarrow \text{anti_discriminative_function}(\tilde{A}(i), \Theta)$
54:	$\text{winner_model} \leftarrow \text{arg_max}(M, \max(\tilde{Y}))$
55:	return winner_model
56:	end procedure

12. sz. melléklet: A genetikai potenciál kereséshez alkalmazott alappopuláció nyersadatai

Tanulóminta redukálása (db)	Objektum azonosító	f ₁	f ₂	f ₃	f ₄	f ₅	f ₆	f ₇	f ₈	f ₉	f ₁₀	Y
0	o1	11.6	8.9	60.9	30.1	53.0	13.0	72.0	12.0	2.5	2.1	548.67
5	o2	12.1	8.8	61.1	30.1	51.0	13.0	68.0	12.0	2.6	2.2	592.27
10	o3	12.3	8.7	61.9	30.1	48.0	13.0	63.0	12.0	2.7	2.2	604.26
5	o4	9.7	6.2	59.9	30.8	52.0	13.0	72.0	10.0	2.5	1.8	567.69
10	o5	9.1	5.6	60.7	30.8	48.0	13.0	68.0	10.0	2.2	1.6	546.26
5	o6	11.8	8.9	61.8	30.0	48.0	13.0	67.0	12.0	2.5	2.1	555.07
10	o7	11.8	9.0	62.8	29.7	44.0	13.0	62.0	12.0	2.4	2.2	542.22
5	o8	10.4	8.4	57.4	29.8	53.0	13.0	72.0	12.0	2.7	2.0	585.15
10	o9	10.3	8.8	52.2	27.2	51.0	13.0	72.0	12.0	3.0	1.8	608.33
5	o10	12.1	8.8	61.1	30.1	51.0	13.0	68.0	12.0	2.6	2.2	594.83
10	o11	12.3	8.7	61.9	30.1	48.0	13.0	63.0	12.0	2.7	2.2	604.26
5	o12	11.3	8.8	62.1	29.7	49.0	13.0	67.0	12.0	2.2	2.0	580.09
10	o13	10.5	7.6	63.3	29.7	48.0	13.0	65.0	11.0	2.0	2.0	518.18
5	o14	12.2	8.8	60.5	30.2	52.0	13.0	69.0	12.0	2.7	2.2	587.23
10	o15	12.4	8.7	61.2	30.2	48.0	13.0	64.0	12.0	2.7	2.2	590.72
5	o16	10.9	8.3	61.6	30.1	49.0	13.0	69.0	11.0	2.3	1.9	576.42
10	o17	10.7	8.3	62.1	30.1	46.0	12.0	65.0	10.0	2.2	1.8	554.11
5	o18	11.8	8.8	61.6	30.1	48.0	13.0	67.0	12.0	2.6	2.1	582.61
10	o19	12.1	8.8	62.5	30.0	43.0	13.0	62.0	12.0	2.6	2.2	578.95
5	o20	11.0	8.4	61.6	30.0	48.0	13.0	68.0	11.0	2.3	1.9	573.92
10	o21	10.4	8.4	60.8	30.7	43.0	13.0	68.0	7.0	2.4	1.7	535.71

Forrás: Saját szerkesztés

13. sz. melléklet: A genetikai potenciál kereséshez alkalmazott alappopuláció rangsorai attribútumonként

Tanulóminta redukálása (db)	Objektum azonosító	f ₁	f ₂	f ₃	f ₄	f ₅	f ₆	f ₇	f ₈	f ₉	f ₁₀	Y
0	o1	11	3	7	12	20	2	18	8	11	11	548.67
5	o2	6	9	8	14	15	2	11	8	7	6	592.27
10	o3	2	12	15	9	5	2	3	8	4	2	604.26
5	o4	20	20	3	21	18	2	18	2	13	19	567.69
10	o5	21	21	5	20	5	2	11	2	19	21	546.26
5	o6	10	2	14	7	5	2	8	8	12	9	555.07
10	o7	9	1	20	4	3	2	1	8	15	8	542.22
5	o8	17	15	2	5	20	2	18	8	3	12	585.15
10	o9	19	7	1	1	15	2	18	8	1	17	608.33
5	o10	6	9	8	14	15	2	11	8	7	6	594.83
10	o11	2	12	15	9	5	2	3	8	4	2	604.26
5	o12	12	5	18	2	13	2	8	8	18	13	580.09
10	o13	16	19	21	3	5	2	6	5	21	14	518.18
5	o14	4	8	4	18	18	2	16	8	6	5	587.23
10	o15	1	11	10	17	5	2	5	8	2	1	590.72
5	o16	14	18	11	13	13	2	16	5	16	15	576.42
10	o17	15	17	17	16	4	1	6	2	20	18	554.11
5	o18	8	4	12	11	5	2	8	8	10	10	582.61
10	o19	5	6	19	6	1	2	1	8	9	4	578.95
5	o20	13	16	13	8	5	2	11	5	17	16	573.92
10	o21	18	14	6	19	1	2	11	1	14	20	535.71

Forrás: Saját szerkesztés

14. sz. melléklet: A genetikai potenciál kereséshez alkalmazott populáció rangsorai a második iterációban

Tanulóminta redukálása (db)	Objektum azonosító	f ₁	f ₂	f ₃	f ₄	f ₅	f ₆	f ₇	f ₈	f ₉	f ₁₀	Y
0	o1	11	3	8	13	21	2	18	8	12	11	548.67
5	o2	6	10	9	15	15	2	11	8	8	6	592.27
10	o3	2	13	16	10	5	2	3	8	5	2	604.26
5	o4	21	21	4	22	19	2	18	2	14	20	567.69
10	o5	22	22	6	21	5	2	11	2	20	22	546.26
5	o6	10	2	15	8	5	2	8	8	13	9	555.07
10	o7	9	1	21	5	3	2	1	8	16	8	542.22
5	o8	18	16	3	6	21	2	18	8	4	12	585.15
10	o9	6	10	9	15	15	2	11	8	8	6	608.33
5	o10	2	13	16	10	5	2	3	8	5	2	594.83
10	o11	12	6	19	3	13	2	8	8	19	13	604.26
5	o12	17	20	22	4	5	2	6	5	22	14	580.09
10	o13	4	9	5	19	19	2	16	8	7	5	518.18
5	o14	1	12	11	18	5	2	5	8	3	1	587.23
10	o15	14	19	12	14	13	2	16	5	17	15	590.72
5	o16	16	18	18	17	4	1	6	2	21	19	576.42
10	o17	8	4	13	12	5	2	8	8	11	10	554.11
5	o18	5	7	20	7	1	2	1	8	10	4	582.61
10	o19	13	17	14	9	5	2	11	5	18	17	578.95
5	o20	19	15	7	20	1	2	11	1	15	21	573.92
10	o21	6	10	9	15	15	2	11	8	8	6	535.71
13	o22	15	5	1	1	15	2	18	8	1	16	619.25

Forrás: Saját szerkesztés

15. sz. melléklet: A genetikai potenciál kereséshez alkalmazott populáció nyersadatai a 11. iterációban

Tanulóminta redukálása (db)	Objektum azonosító	f ₁	f ₂	f ₃	f ₄	f ₅	f ₆	f ₇	f ₈	f ₉	f ₁₀	Y
0	o1	11.6	8.9	60.9	30.1	53.0	13.0	72.0	12.0	2.5	2.1	548.67
5	o2	12.1	8.8	61.1	30.1	51.0	13.0	68.0	12.0	2.6	2.2	592.27
10	o3	12.3	8.7	61.9	30.1	48.0	13.0	63.0	12.0	2.7	2.2	604.26
5	o4	9.7	6.2	59.9	30.8	52.0	13.0	72.0	10.0	2.5	1.8	567.69
10	o5	9.1	5.6	60.7	30.8	48.0	13.0	68.0	10.0	2.2	1.6	546.26
5	o6	11.8	8.9	61.8	30.0	48.0	13.0	67.0	12.0	2.5	2.1	555.07
10	o7	11.8	9.0	62.8	29.7	44.0	13.0	62.0	12.0	2.4	2.2	542.22
5	o8	10.4	8.4	57.4	29.8	53.0	13.0	72.0	12.0	2.7	2.0	585.15
10	o9	10.3	8.8	52.2	27.2	51.0	13.0	72.0	12.0	3.0	1.8	608.33
5	o10	12.1	8.8	61.1	30.1	51.0	13.0	68.0	12.0	2.6	2.2	594.83
10	o11	12.3	8.7	61.9	30.1	48.0	13.0	63.0	12.0	2.7	2.2	604.26
5	o12	11.3	8.8	62.1	29.7	49.0	13.0	67.0	12.0	2.2	2.0	580.09
10	o13	10.5	7.6	63.3	29.7	48.0	13.0	65.0	11.0	2.0	2.0	518.18
5	o14	12.2	8.8	60.5	30.2	52.0	13.0	69.0	12.0	2.7	2.2	587.23
10	o15	12.4	8.7	61.2	30.2	48.0	13.0	64.0	12.0	2.7	2.2	590.72
5	o16	10.9	8.3	61.6	30.1	49.0	13.0	69.0	11.0	2.3	1.9	576.42
10	o17	10.7	8.3	62.1	30.1	46.0	12.0	65.0	10.0	2.2	1.8	554.11
5	o18	11.8	8.8	61.6	30.1	48.0	13.0	67.0	12.0	2.6	2.1	582.61
10	o19	12.1	8.8	62.5	30.0	43.0	13.0	62.0	12.0	2.6	2.2	578.95
5	o20	11.0	8.4	61.6	30.0	48.0	13.0	68.0	11.0	2.3	1.9	573.92
10	o21	10.4	8.4	60.8	30.7	43.0	13.0	68.0	7.0	2.4	1.7	535.71
13	o22	10.8	8.8	49.5	26.0	51.0	13.0	72.0	12.0	3.1	1.9	619.25
14	o23	10.9	8.9	48.8	25.8	50.0	13.0	72.0	11.0	3.2	1.9	627.62
15	o24	11.0	9.0	48.2	25.7	49.0	13.0	72.0	10.0	3.2	2.0	605.81
27	o25	11.9	8.7	50.7	26.2	43.0	12.0	59.0	11.0	3.4	2.1	611.11
31	o26	12.1	8.7	51.4	26.3	40.0	11.0	55.0	11.0	3.4	2.2	625.51
16	o27	10.9	9.1	47.6	25.6	49.0	13.0	72.0	10.0	3.3	1.9	599.16
32	o28	12.3	8.7	50.6	26.1	40.0	11.0	55.0	11.0	3.5	2.2	623.48
29	o29	12.1	8.7	51.2	26.3	40.0	12.0	56.0	11.0	3.5	2.1	610.44
31	o30	12.4	8.7	50.8	26.1	40.0	11.0	71.0	10.0	3.5	2.2	623.48
13	o31	10.0	9.0	49.2	26.5	48.0	13.0	72.0	9.0	3.2	1.8	635.98
14	o32	10.0	9.0	49.4	26.4	48.0	12.0	71.0	9.0	3.2	1.8	636.73

Forrás: Saját szerkesztés

16. sz. melléklet: A genetikai potenciál kereséshez alkalmazott populáció rangsorszámai a 11. iterációban

Tanulóminta redukálása (db)	Objektum azonosító	f ₁	f ₂	f ₃	f ₄	f ₅	f ₆	f ₇	f ₈	f ₉	f ₁₀	Y
0	o1	17	7	18	23	31	8	1	18	22	18	548.67
5	o2	22	15	19	25	25	8	14	18	18	24	592.27
10	o3	28	22	26	20	10	8	25	18	15	29	604.26
5	o4	2	31	14	32	29	8	1	4	24	3	567.69
10	o5	1	32	16	31	10	8	14	4	30	1	546.26
5	o6	18	6	25	18	10	8	19	18	23	21	555.07
10	o7	19	4	31	15	8	8	27	18	26	22	542.22
5	o8	7	26	13	16	31	8	1	18	14	16	585.15
10	o9	5	13	12	12	25	8	1	18	12	5	608.33
5	o10	22	15	19	25	25	8	14	18	18	24	594.83
10	o11	28	22	26	20	10	8	25	18	15	29	604.26
5	o12	16	11	29	13	20	8	19	18	29	15	580.09
10	o13	8	30	32	14	10	8	22	10	32	13	518.18
5	o14	27	14	15	29	29	8	12	18	17	26	587.23
10	o15	32	20	21	28	10	8	24	18	13	31	590.72
5	o16	13	29	22	24	20	8	12	10	27	12	576.42
10	o17	9	28	28	27	9	4	22	4	31	4	554.11
5	o18	20	9	23	22	10	8	19	18	21	20	582.61
10	o19	25	12	30	17	5	8	27	18	20	27	578.95
5	o20	15	27	24	19	10	8	14	10	28	8	573.92
10	o21	6	25	17	30	5	8	14	1	25	2	535.71
13	o22	10	10	6	4	25	8	1	18	11	11	619.25
14	o23	11	8	3	3	24	8	1	10	10	10	627.62
15	o24	14	5	2	2	20	8	1	4	7	14	605.81
27	o25	21	21	8	7	5	4	29	10	5	17	611.11
31	o26	26	19	11	9	1	1	31	10	4	23	625.51
16	o27	12	1	1	1	20	8	1	4	6	9	599.16
32	o28	30	18	7	5	1	1	31	10	1	28	623.48
29	o29	24	24	10	8	1	4	30	10	3	19	610.44
31	o30	31	17	9	6	1	1	10	4	2	32	623.48
13	o31	3	3	4	11	10	8	1	2	9	7	635.98
14	o32	4	2	5	10	10	4	10	2	8	6	636.73

Forrás: Saját szerkesztés

17. sz. melléklet: Felügyelet nélküli modellek objektumleíró tulajdonságai

Attribútum	Leírás	Értelmezés
<i>T01: Válaszképesség</i>	Bináris változó, azt mutatja meg, hogy adott modell adott objektumra (auditjelentésre) vonatkozóan képes volt-e gyanút megállapítani.	Értéke 0, abban az esetben, ha a modell nem volt képes adott objektum tekintetében gyanú rögzítésére, ellenkezőleg az értéke 1. A modellektől azt várjuk el alapesetben, hogy képesek legyenek választ adni, azaz gyanúba hozni egy vagy több kontrollt. Amelyik modell nem válaszképessé, az feltételezhetően kevésbé alkalmas a problémára megoldást találni. Kontextus független attribútum. Írány-preferencia: minél nagyobb, annál jobb.
<i>T02: Összes rendszerválasz</i>	A modellek a vélt gyanús objektumok egy halmazával térnek vissza, tehát az attribútum a rendszerválaszok összegét regisztrálja.	Értéke 0 és a kontrollok maximuma * k között mozog. Minél kevesebb a rendszerválasz, tehát a gyanús objektumok száma, feltételezhetően, annál homogénebb döntésre jutott a modell, annál kisebb a véletlenszerűség kockázata. Kontextus független attribútum. Írány-preferencia: minél kisebb, annál jobb.
<i>T03: Egyedi rendszerválaszok száma</i>	Az összes egyedi rendszerválasz redundancia nélkül.	Értéke 1 és a kontrollok maximuma között mozog, de maximum annyi, mint T02. Hasonlóan a T02 attribútumhoz, akkor preferált egy modell, ha kevesebb egyedi rendszerválasszal tér vissza. Kontextus független attribútum. Írány-preferencia: minél kisebb, annál jobb.
<i>T04: Egyedüli rendszerválaszok száma</i>	Azon gyanús objektumok halmaza, melyeket egy adott modell kizárólag egyedüli gyanúként kezel a kiugró esetek speciális alakzataként, mely vezethet rendszeranomáliára.	Értéke 0 és a T02 értéke között mozog. Minél kevesebb az egyedi eset, vélhetően annál kevesebb az előforduló rendellenesség. Kontextus független attribútum. Írány-preferencia: minél kisebb, annál jobb.
<i>T05: Összes rendszerválaszra jutó egyedi rendszerválaszok száma</i>	Származtatott attribútum, a T03 és a T02 hányadosa.	Értéke 0 és 1 között mozog. Optimummal rendelkező attribútum, tehát az irány-preferencia meghatározása nem triviális. Kontextus független attribútum. Írány-preferencia: Feltételezhetően minél nagyobb, annál jobb, azonban ez további alátámasztást igényel.
<i>T06: Egyedi rendszerválaszokra jutó egyedüli rendszerválaszok száma</i>	Származtatott attribútum, a T04 és a T03 hányadosa.	Értéke 0 és 1 között mozog. Optimummal rendelkező attribútum, tehát az irány-preferencia meghatározása nem triviális. Kontextus független attribútum. Írány-preferencia: Feltételezhetően minél nagyobb, annál jobb, azonban ez további alátámasztást igényel.
<i>T07: Elemi tömörsűrűségek száma</i>	A rendszerválaszok közötti objektumonként elkülöníthető gyanúmomentumok tömörszerű szakadásainak száma.	Értéke 0 és T02 között mozog. A rendszer a generált gyanús outputokat rendezetten a leghasonlóbb objektum szerint csökkenő sorrendben nyújtja vissza a döntéshozó kezébe, ezért előfordulhat, hogy egy adott generált gyanú tömbösítve jelenik meg. Egy adott gyanú esetén ezért a tömbösítés jelentheti annak erőteljesebb súlyozását, tehát a tömörsűrűség jelei előállhatnak rendszerzavar vagy véletlenszerűség okán, azaz anomáliaként. A mutató ezen

		tömbszakadások számát méri. Kontextus független attribútum.
		Írány-preferencia: minél kisebb, annál jobb.
<i>T08: Elemi homogenitás mutató</i>	A rendszerválaszok homogenitását méri páronkénti összehasonlítás által a logikai ekvivalencia művelet felhasználásával. Minden ekvivalens rendszerválasz esetén az érték 1, különben 0. A mutató ezen értékeket összesíti és az összes összehasonlítással elosztja.	Értéke 0 és 1 között mozog. A homogénebb rendszerválaszok egyfajta önerősítő mechanizmusok. A heterogén rendszerválaszok számottevő gyanút jelölhetnek meg függetlenül, mely nagyobb véletlenszerűségi kockázatot jelent. Kontextus független attribútum.
		Írány-preferencia: minél nagyobb, annál jobb.
<i>T09: Egyedi mintázati rendszerválaszok száma</i>	A gyanúként megjelölt kontrollok objektumonként megjelenő rendszerválaszainak száma, azaz a top k auditjelentés közül hány esetben volt ugyanazon megállapítások halmaza gyanús ítélet tömbszerűen. Hasonlóan a homogenitást erősíti.	Értéke 1 a kiválasztott k érték között mozog. Kontextus független attribútum.
		Írány-preferencia: minél kisebb, annál jobb.
<i>T10: Egyedüli mintázati rendszerválaszok száma</i>	Azon gyanús objektumok tömbösített halmaza, melyeket egy adott modell kizárólag egyedüli gyanúként kezel a kiugró esetek speciális alakzataként, mely vezethet rendszeranomáliára.	Értéke 0 és T09 között mozog. Minél kevesebb az egyedi mintaszerűen megjelenő eset, vélhetően annál kevesebb az előforduló rendellenesség. Kontextus független attribútum.
		Írány-preferencia: minél kisebb, annál jobb.
<i>T11: Összes rendszerválaszra jutó egyedi mintázati rendszerválaszok száma</i>	Származtatott attribútum, a T09 és a k hányadosa.	Értéke 0 és 1 között mozog. Optimummal rendelkező attribútum, tehát az irány-preferencia meghatározása nem triviális. Kontextus független attribútum.
		Írány-preferencia: Feltételezhetően minél nagyobb, annál jobb, azonban ez további alátámasztást igényel.
<i>T12: Egyedi mintázati rendszerválaszokra jutó egyedüli mintázati rendszerválaszok száma</i>	Származtatott attribútum, a T09 és a T10 hányadosa.	Értéke 0 és 1 között mozog. Optimummal rendelkező attribútum, tehát az irány-preferencia meghatározása nem triviális. Kontextus független attribútum.
		Írány-preferencia: Feltételezhetően minél nagyobb, annál jobb, azonban ez további alátámasztást igényel.
<i>T13: Mintázati tömbösítések száma</i>	Mintázati gyanúk tömbösítéseiinek száma, hasonlóan értelemezendő, mint a T07 attribútum.	Értéke 0 és k között mozog. A rendszer a generált gyanús outputokat rendezetten a leghasonlóbb objektum szerint csökkenő sorrendben nyújtja vissza a döntéshozó kezébe, ezért előfordulhat, hogy egy adott generált gyanú tömbösítve jelenik meg. Egy adott gyanú-mintázat esetén ezért a tömbösítés jelentheti annak erőteljesebb súlyozását, tehát a tömbösítés jelei előállhatnak rendszerzavar vagy véletlenszerűség okán, azaz anomáliaként. A mutató ezen tömbszakadások számát méri, ahol egy minta az objektum gyanús kontrolljait jelöli. Kontextus független attribútum.
		Írány-preferencia: minél kisebb, annál jobb.

<i>T14: Mintázati homogenitás mutató</i>	A rendszerválaszok mintázatszerű homogenitását méri páronkénti összehasonlítás által a logikai ekvivalencia művelet felhasználásával.	Értéke 0 és 1 között mozog. A homogénebb rendszerválaszok egyfajta önerősítő mechanizmusok. A heterogén rendszerválaszok számottevő gyanút jelölhetnek meg függetlenül, mely nagyobb véletlenszerűségi kockázatot jelent. Kontextus független attribútum. Írány-preferencia: minél nagyobb, annál jobb.
<i>T15: Összes egyedi csoportszám</i>	Az összes egyedi egy azon kontrollterülethez tartozó gyanús kontrollok csoportszáma az ISO IEC 27001:2013 A mellékletének csoportosítása szerint, mely tartalmazza a „nincs megállapítás” halmazát, továbbá az egy kontrollterülettel kiegészített besorolást.	Értéke 0 és 16 között mozog. Minél kevesebb az összes egyedi csoportszám, annál homogénebb a rendszerválasz csökkentve a véletlenszerűségi kockázatot. Kontextusfüggő attribútum Írány-preferencia: minél kisebb, annál jobb.
<i>T16: Csoport alapú tömörsérülések száma</i>	Csoportok (kontrollterület) alapú gyanúk tömörsérüléseinek száma, hasonlóan értelmezendő, mint a T07 attribútum.	Értéke 0 és T02 között mozog. A rendszer a generált gyanús outputokat rendezetten a leghasonlóbb objektum szerint csökkenő sorrendben nyújtja vissza a döntéshozó kezébe, ezért előfordulhat, hogy egy adott generált gyanú tömbösítve jelenik meg. Egy adott gyanú-csoport esetén ezért a tömbösítés jelentheti annak erőteljesebb súlyozását, tehát a tömörsérülés jelei előállhatnak rendszerzavar vagy véletlenszerűség okán, azaz anomáliaként. A mutató ezen tömörszakadások számát méri, ahol egy csoport az objektum gyanús kontrolljait jelöli az ISO IEC 27001:2013 A mellékletében meghatározott kontrollterületek alapján. Kontextus függő attribútum. Írány-preferencia: minél kisebb, annál jobb.
<i>T17: Csoport alapú homogenitás mutató</i>	A rendszerválaszok csoport alapú homogenitását méri páronkénti összehasonlítás által a logikai ekvivalencia művelet felhasználásával.	Értéke 0 és 1 között mozog. A homogénebb rendszerválaszok egyfajta önerősítő mechanizmusok. A heterogén rendszerválaszok számottevő gyanút jelölhetnek meg függetlenül, mely nagyobb véletlenszerűségi kockázatot jelent. Kontextus függő attribútum. Írány-preferencia: minél nagyobb, annál jobb.

Forrás: Saját szerkesztés

18. sz. melléklet: Felügyelet nélküli modellek objektumleíró tulajdonságainak irány-preferenciát ellenőrző korrelációs mátrixa

Attribútum azonosító		T01	T02	T03	T04	T05	T06	T07	T08	T09	T10	T11	T12	T13	T14	T15	T16	T17
	<i>irány-preferencia</i>	0	1	1	1	0	0	1	0	1	1	0	0	1	0	1	1	0
T01	0	1.00																
T02	1	-0.66	1.00															
T03	1	-0.55	0.97	1.00														
T04	1	-0.62	0.88	0.93	1.00													
T05	0	0.83	-0.41	-0.21	-0.27	1.00												
T06	0	-0.45	0.57	0.64	0.80	-0.11	1.00											
T07	1	-0.59	0.91	0.92	0.90	-0.29	0.63	1.00										
T08	0	0.10	-0.71	-0.81	-0.66	-0.23	-0.51	-0.73	1.00									
T09	1	-0.36	0.87	0.93	0.81	0.00	0.58	0.84	-0.90	1.00								
T10	1	-0.47	0.89	0.94	0.88	-0.09	0.69	0.86	-0.83	0.98	1.00							
T11	0	0.70	-0.11	0.05	-0.13	0.86	-0.18	-0.06	-0.38	0.29	0.17	1.00						
T12	0	-0.26	0.68	0.77	0.75	0.15	0.82	0.71	-0.81	0.85	0.89	0.25	1.00					
T13	1	-0.76	0.69	0.65	0.68	-0.56	0.59	0.74	-0.48	0.54	0.58	-0.52	0.49	1.00				
T14	0	0.16	-0.76	-0.84	-0.68	-0.15	-0.46	-0.75	0.96	-0.95	-0.88	-0.39	-0.78	-0.47	1.00			
T15	1	-0.46	0.92	0.94	0.84	-0.17	0.65	0.85	-0.84	0.94	0.93	0.09	0.82	0.64	-0.87	1.00		
T16	1	-0.63	0.83	0.81	0.74	-0.35	0.56	0.87	-0.74	0.79	0.79	-0.17	0.69	0.86	-0.74	0.80	1.00	
T17	0	0.13	-0.71	-0.78	-0.67	-0.11	-0.55	-0.73	0.95	-0.82	-0.77	-0.23	-0.76	-0.57	0.89	-0.84	-0.74	1.00

Forrás: Saját szerkesztés

19. sz. melléklet: Felügyelet nélküli modellek leíró statisztikái

Modell azonosító	Rangsor átlag	Rangsor maximum	Rangsor miniumum	Rangsor szórás	Percentilis- rang átlag	Percentilis- rang maximum	Percentilis- rang minimum	Percentilis- rank szórás	Norm. értékek átlag	Norm. értékek maximum	Norm. értékek minimum	Norm. értékek szórás
<i>R_NS_NRS_C</i>	63.5784	108.2941	13.8235	37.6245	0.4883	0.8363	0.1610	0.2731	0.4078	0.8235	0.0000	0.3350
<i>R_NS_NRS_E</i>	43.7451	62.1765	17.7059	12.9052	0.6460	0.8215	0.5271	0.0859	0.6059	0.7828	0.4836	0.0798
<i>R_NS_NRS_P</i>	61.7304	108.2941	13.8235	37.6709	0.5016	0.8363	0.1610	0.2736	0.4232	0.8235	0.0000	0.3355
<i>R_NS_RS_C</i>	63.5784	108.2941	13.8235	37.6245	0.4883	0.8363	0.1610	0.2731	0.4078	0.8235	0.0000	0.3350
<i>R_NS_RS_E</i>	55.3431	74.5294	17.7059	18.0978	0.5566	0.8215	0.3975	0.1355	0.5009	0.7828	0.3305	0.1576
<i>R_NS_RS_P</i>	61.7549	108.2941	13.8235	37.6694	0.5014	0.8363	0.1610	0.2736	0.4229	0.8235	0.0000	0.3355
<i>R_S_NRS_C</i>	65.1716	108.2941	13.8235	37.4382	0.4775	0.8363	0.1610	0.2715	0.3962	0.8235	0.0000	0.3328
<i>R_S_NRS_E</i>	44.1422	62.2353	17.7059	13.2433	0.6432	0.8215	0.5267	0.0883	0.6024	0.7828	0.4795	0.0835
<i>R_S_NRS_P</i>	63.5441	108.2941	13.8235	37.3630	0.4899	0.8363	0.1610	0.2711	0.4125	0.8235	0.0000	0.3326
<i>R_S_RS_C</i>	68.8284	108.2941	13.8235	33.6457	0.4539	0.8363	0.1610	0.2472	0.3697	0.8050	0.0000	0.3029
<i>R_S_RS_E</i>	60.2304	74.5294	17.7059	16.2520	0.5178	0.8215	0.3975	0.1181	0.4583	0.7828	0.3305	0.1343
<i>R_S_RS_P</i>	66.5784	108.2941	13.8235	34.0070	0.4693	0.8363	0.1610	0.2498	0.3864	0.8050	0.0000	0.3064

Forrás: Saját szerkesztés

20. sz. melléklet: Felügyelet nélküli modellek szórás irány-preferenciáinak értékelése varianciaelemzéssel

<i>Varianciaelemzésben felhasznált és kiértékelt változók</i>	<i>Szórás irány-preferencia</i>	
<i>Irány</i>	111	000
<i>Felhasznált eset</i>	4	3
<i>Átlag</i>	1005.30	1011.67
<i>Szórás</i>	5.46	8.03
<i>Levene-teszt</i>	0.36	
<i>F-próba</i>	1.54	
<i>Szignifikancia</i>	0.27	

Forrás: Saját szerkesztés

21. sz. melléklet: Modell-leíró tulajdonságokra illesztett anti-diszkriminatív függvény súlysámai (első irány-preferencia nézet)

Lépcsők	x₁	x₂	x₃	x₄	x₅	x₆	x₇	x₈	x₉	x₁₀	x₁₁	x₁₂
S1	904.4	11	11	21.5	11	11	11	11	11	44.1	11	11
S2	903.4	10	10	20.5	10	10	10	10	10	43.1	10	10
S3	902.4	9	9	19.5	9	9	9	9	9	42.1	9	9
S4	901.4	8	8	18.5	8	8	8	8	8	41.1	8	8
S5	900.4	7	7	17.5	7	7	7	7	7	40.1	7	7
S6	899.3	6	6	16.5	6	6	6	6	6	39.1	6	6
S7	893.8	5	5	6.5	5	5	5	5	5	38.1	5	5
S8	892.8	4	4	5.5	4	4	4	4	4	4	4	4
S9	891.8	3	3	3	3	3	3	3	3	3	3	3
S10	890.8	2	2	2	2	2	2	2	2	2	2	2
S11	889.8	1	1	1	1	1	1	1	1	1	1	1
S12	888.8	0	0	0	0	0	0	0	0	0	0	0

Forrás: Saját szerkesztés

22. sz. melléklet: Felügyelet nélküli modellek leíró tulajdonságainak értékelése varianciaelemzéssel (első irány-preferencia nézet)

Varianciaelemzésben felhasznált és kiértékelt változók	Megállapítások súlyozása szerinti kategóriák			
	Naiv átlagolás		Hasonlóságelemzés	
<i>Idealitás mutató</i>				
<i>Súlyozás</i>	Nincs súlyozás	Van súlyozás	Nincs súlyozás	Van súlyozás
<i>Felhasznált eset</i>	6	6	6	6
<i>Átlag</i>	5.13	5.33	1000.17	1000.33
<i>Szórás</i>	1.12	1.01	7.08	5.96
<i>Levene-teszt</i>	1.00		0.79	
<i>F-próba</i>	0.12		0.02	
<i>Szignifikancia</i>	0.74		0.97	

Varianciaelemzésben felhasznált és kiértékelt változók	Rendszerek súlyozása szerinti kategóriák			
	Naiv átlagolás		Hasonlóságelemzés	
<i>Idealitás mutató</i>				
<i>Súlyozás</i>	Nincs súlyozás	Van súlyozás	Nincs súlyozás	Van súlyozás
<i>Felhasznált eset</i>	6	6	6	6
<i>Átlag</i>	4.88	5.58	1003.00	997.50
<i>Szórás</i>	1.23	0.68	6.20	5.39
<i>Levene-teszt</i>	0.12		0.92	
<i>F-próba</i>	1.52		2.69	
<i>Szignifikancia</i>	0.25		0.13	

Varianciaelemzésben felhasznált és kiértékelt változók	Távolság/kapcsolat-metrikák szerinti kategóriák					
	Naiv átlagolás			Hasonlóságelemzés		
<i>Idealitás mutató</i>						
<i>Távolság/kapcsolat metrikák</i>	EUC	PEA	COS	EUC	PEA	COS
<i>Felhasznált eset</i>	4	4	4	4	4	4
<i>Átlag</i>	4.06	5.69	5.94	999.25	1003.50	998.00
<i>Szórás</i>	0.85	0.52	0.24	10.21	3.00	2.71
<i>Levene-teszt</i>	0.07			0.01		
<i>F-próba</i>	11.88			0.83		
<i>Szignifikancia</i>	0.00			0.47		

Forrás: Saját szerkesztés

23. sz. melléklet: Modell-leíró tulajdonságokra illesztett anti-diszkriminatív függvény súlyszámai
(második irány-preferencia nézet)

Lépcsők	X1	X2	X3	X4	X5	X6	X7	X8	X9	X10	X11	X12
S1	490.8	11	11	924.1	11	11	11	12.5	11	11	11	11
S2	489.8	10	10	923.1	10	10	10	11.5	10	10	10	10
S3	488.8	9	9	922.1	9	9	9	10.5	9	9	9	9
S4	487.8	8	8	921.1	8	8	8	9.5	8	8	8	8
S5	7	7	7	920.1	7	7	7	8.5	7	7	7	7
S6	6	6	6	919.1	6	6	6	7.5	6	6	6	6
S7	5	5	5	918.1	5	5	5	6.5	5	5	5	5
S8	4	4	4	917.1	4	4	4	5.5	4	4	4	4
S9	3	3	3	451.3	3	3	3	4.5	3	3	3	3
S10	2	2	2	450.3	2	2	2	3.5	2	2	2	2
S11	1	1	1	446.3	1	1	1	2.5	1	1	1	1
S12	0	0	0	443.8	0	0	0	0	0	0	0	0

Forrás: Saját szerkesztés

24. sz. melléklet: Felügyelet nélküli modellek leíró tulajdonságainak értékelése varianciaelemzéssel (második irány-preferencia nézet)

Varianciaelemzésben felhasznált és kiértékelt változók	Megállapítások súlyozása szerinti kategóriák			
	Naiv átlagolás		Hasonlóságelemzés	
<i>Idealitás mutató</i>				
Súlyozás	Nincs súlyozás	Van súlyozás	Nincs súlyozás	Van súlyozás
<i>Felhasznált eset</i>	6	6	6	6
<i>Átlag</i>	4.46	6.00	1009.50	1000.50
<i>Szórás</i>	1.20	0.85	9.05	10.86
<i>Levene-teszt</i>	0.38		0.48	
<i>F-próba</i>	6.60		9.73	
<i>Szignifikancia</i>	0.03		0.01	

Varianciaelemzésben felhasznált és kiértékelt változók	Rendszerek súlyozása szerinti kategóriák			
	Naiv átlagolás		Hasonlóságelemzés	
<i>Idealitás mutató</i>				
Súlyozás	Nincs súlyozás	Van súlyozás	Nincs súlyozás	Van súlyozás
<i>Felhasznált eset</i>	6	6	6	6
<i>Átlag</i>	4.83	5.63	1004.50	996.50
<i>Szórás</i>	1.01	1.41	9.93	16.03
<i>Levene-teszt</i>	0.37		0.24	
<i>F-próba</i>	1.17		1.08	
<i>Szignifikancia</i>	0.31		0.32	

Varianciaelemzésben felhasznált és kiértékelt változók	Távolság/kapcsolat-metrikák szerinti kategóriák					
	Naiv átlagolás			Hasonlóságelemzés		
<i>Idealitás mutató</i>						
<i>Távolság/kapcsolat metrikák</i>	EUC	PEA	COS	EUC	PEA	COS
<i>Felhasznált eset</i>	4	4	4	4	4	4
<i>Átlag</i>	6.06	4.44	5.19	999.00	1005.75	996.75
<i>Szórás</i>	0.31	1.56	1.30	2.45	18.71	15.56
<i>Levene-teszt</i>	0.12			0.08		
<i>F-próba</i>	1.88			0.44		
<i>Szignifikancia</i>	0.21			0.66		

Forrás: Saját szerkesztés

25. sz. melléklet: Felügyelet nélküli modellek kategóriáinak értékelése varianciaelemzéssel

Varianciaelemzésben felhasznált és kiértékelt változók	Megállapítások súlyozása szerinti kategóriák	
<i>Idealitás mutató</i>	Hasonlóságelemzés	
<i>Kategória</i>	Nincs súlyozás	Van súlyozás
<i>Felhasznált eset</i>	6	6
<i>Átlag</i>	1003.77	996.27
<i>Szórás</i>	7.80	9.50
<i>Levene-teszt szignifikanciája</i>		0.69
<i>F-próba</i>		2.32
<i>Szignifikancia</i>		0.17

Varianciaelemzésben felhasznált és kiértékelt változók	Rendszerek súlyozása szerinti kategóriák	
<i>Idealitás mutató</i>	Hasonlóságelemzés	
<i>Kategória</i>	Nincs súlyozás	Van súlyozás
<i>Felhasznált eset</i>	6	6
<i>Átlag</i>	1002.27	997.77
<i>Szórás</i>	11.09	997.77
<i>Levene-teszt szignifikanciája</i>		0.17
<i>F-próba</i>		0.70
<i>Szignifikancia</i>		0.42

Varianciaelemzésben felhasznált és kiértékelt változók	Távolság/kapcsolat-metrikák szerinti kategóriák		
<i>Idealitás mutató</i>	Hasonlóságelemzés		
<i>Kategória</i>	EUC	PEA	COS
<i>Felhasznált eset</i>	4	4	4
<i>Átlag</i>	1007.60	1001.60	990.85
<i>Szórás</i>	8.04	4.62	5.50
<i>Levene-teszt szignifikanciája</i>			0.15
<i>F-próba</i>			7.43
<i>Szignifikancia</i>			0.01
<i>Scheffé-próba szignifikanciája (EUC – PEA)</i>			0.43
<i>Scheffé-próba szignifikanciája (EUC – COS)</i>			0.01
<i>Scheffé-próba szignifikanciája (PEA – COS)</i>			0.10

Forrás: Saját szerkesztés

Köszönetnyilvánítás

Ezúton szeretném megköszönni témavezetőmnek, **Dr. Pitlik Lászlónak**, a kutatás és dolgozatírás során nyújtott folyamatos támogatását, szakmai segítségét és hasznos tanácsait, és azt, hogy bevezetett a mesterséges intelligencia világába.

Köszönettel tartozom opponenseimnek, **Dr. Bánkuti Gyöngyi Ilonának**, **Dr. Kása Richárdnak** és **Dr. Szalay Zsigmond Gábornak** a dolgozat minőségének javítása érdekében tett értékes és konstruktív észrevételeikért.

Hálás köszönettel tartozom **családomnak**, **barátaimnak**, **munkatársaimnak**, akik folyamatosan bíztattak, érdeklődtek a kutatás felől és türelmesek voltak hozzám a nehéz időszakban is.

Külön köszönettel tartozom Édesanyámnak, **Sziksza Katalinnak**, hogy mindvégig bízott bennem, nélküle ez a dolgozat nem készülhetett volna el.

Végezetül hálásan köszönöm feleségemnek, **Dr. Marta Léték-nek**, hogy biztatásával tartotta bennem a lelket, mely hatalmas erőt adott a legnehezebb pillanatokban is.