

VILNIAUS UNIVERSITETAS
KAUNO FAKULTETAS

SOCIALINIŲ MOKSLŲ IR TAIKOMOSIOS INFORMATIKOS INSTITUTAS

Finansų technologijų studijų programa
Kodas 6211BX022

MARTYNAS BARTNYKAS

MAGISTRO BAIGIAMASIS DARBAS

**MOKĖJIMŲ STEBĖSENOS METODAS TAIKANTIS PROCESŲ GAVYBĄ IR DIDELIUS
KALBOS MODELIUS**

Kaunas, 2025

VILNIAUS UNIVERSITETAS
KAUNO FAKULTETAS

SOCIALINIŲ MOKSLŲ IR TAIKOMOSIOS INFORMATIKOS INSTITUTAS

MARTYNAS BARTNYKAS

MAGISTRO BAIGIAMASIS DARBAS

MOKĖJIMŲ STEBĖSENOS METODAS TAIKANTIS PROCESŲ GAVYBĄ IR DIDELIUS
KALBOS MODELIUS

Darbo vadovas _____
(parašas)

Partn. Prof. Dr. Darius Dilijonas

(darbo vadovo mokslo laipsnis,
mokslo pedagoginis vardas,
vardas ir pavardė)

Magistrantas _____
(parašas)

Darbo įteikimo data 2025-12-23

Registracijos Nr. _____

Kaunas, 2025

TURINYS

TURINYS.....	3
SANTRUMPŲ SĄRAŠAS	6
SUMMARY	7
IVADAS.....	9
1. PINIGŲ PLOVIMO PREVENCIJOS EFEKTYVUMO DIDINIMO TEORINIAI ASPEKTAI MOKĖJIMŲ STEBĖSENOS KONTEKSTE	12
1.1 Pinigų plovimo prevencijos samprata ir principai.....	12
1.2 Mokėjimų stebėsenos principai ir efektyvumo problematika	16
1.3 Finansų proceso gavybos taikymas mokėjimų stebėsenos sistemose	20
1.4 Didelių kalbos modelių (LLMs) vaidmuo ir pritaikomumas mokėjimų stebėsenoje.....	26
2. SIŪLOMO METODO SKYRIUS	36
2.1 Siūlomo metodo konceptas ir pagrindumas.....	36
2.2 Bendroji metodo struktūra.....	37
2.2.1 Duomenų rinkinio paruošimas	37
2.2.2 Procesų gavybos (Fuzzy Miner) taikymas	39
2.2.3 Didelio kalbos modelio taikymas.....	44
2.2.4 Metodo veikimo schema ir principai.....	54
3. EKSPERIMENTINIS SKYRIUS	59
3.1 Siūlomo hibridinio metodo vykdymas	59
3.1.1. Duomenų rinkinys ir išankstinis duomenų paruošimas.....	59
3.1.2. Procesų gavybos modulio taikymas Celonis platformoje	65
3.1.3. Didelio kalbos modelio (gpt-4.1-mini) modulio taikymas naudojant API sąsają	71
3.2 Rezultatų analizė naudojant H2O AutoML įrankį	74
IŠVADOS.....	86
PASIŪLYMAI	90

LITERATŪROS SĀRAŠAS.....	91
1 PRIEDAS	95
2 PRIEDAS	96
3 PRIEDAS	97
5 PRIEDAS	99
6 PRIEDAS	100
7 PRIEDAS	101
8 PRIEDAS	102

LENTELIŲ IR PAVEIKSLŲ SĄRAŠAS

Paveikslų sąrašas

pav. 1 Pinigų plovimo etapų schema.....	13
pav. 2 Procesų gavybos etapų schema	21
pav. 3 Petri tinklo pavyzdys	22
pav. 4 Didelio kalbos modelio (LLM) veikimo etapų schema.....	28
pav. 5 Išankstinio duomenų apdorojimo įtaka procesų gavybos modelio sudarymui.....	38
pav. 6 Įeinančių briaunų į mazgą A filtravimo procesas Fuzzy Miner algoritme	41
pav. 7 Procesų modelio pavyzdys netaikant ribinių reikšmių filtravimo (kairėje) ir taikant ribinių kraštinių filtravimą (dešinėje)	42
pav. 8 Siūlomo metodo procesų gavybos modulio veikimo schema	43
pav. 9 Užklausų inžinerijos komponentų sąveikos schema	46
pav. 10 Zero-Shot užklausos metodo veikimo schema	47
pav. 11 Siūlomo metodo didelio kalbos modelio veiksmų schema	49
pav. 12 Bendra siūlomo metodo taikymo schema.....	55
pav. 13 Procesų modelis (Process Explorer) sudarytas pagal 1 mėnesio intervalą.....	67
pav. 14 Procesų modelio atvaizdavimas ir analizė naudojant atvejo naršyklės (Case explorer) langą	68

Lentelių sąrašas

1 lentelė Empiriniai duomenys gauti naudojant didelių kalbos modelių technologiją finansinių nusikaltimų ir prevencijos tyrimų srityje	32
2 lentelė Sociodemografiniu kliento profilio šablonu paremta užklausos struktūra	51
3 lentelė Kliento finansinės elgsenos profilio šablonu paremta užklausos struktūra.....	52
4 lentelė Pirminių kintamųjų pavadinimų transformacijos ir reikšmės.....	60
5 lentelė Naujų papildomų kintamųjų transformacijos	63
6 lentelė Procesų gavybos metu išvestos savybės.....	69
7 lentelė Didelio kalbos modelio mokėjimų įverčių savybės.....	73
8 lentelė Naudojamo DRF modelio pirminių rezultatų analizė taikant skirtingus parametrus	76
9 lentelė Naudojamo GBM modelio pirminių rezultatų analizė taikant skirtingus parametrus.....	78
10 lentelė Taikomo metodo metu išgautų savybių galutiniai efektyvumo rezultatai.....	82
11 lentelė Taikomo metodo efektyvumo palyginimas su kitais mokslinėje literatūroje siūlomais metodais	84

SANTRUMPŲ SĄRAŠAS

.
AML – Anti-money laundering (liet. Pinigų plovimo prevencija)
API – Application programming interface (liet. Aplikacijų programavimo sąsaja)
BDAR – Bendrasis duomenų apsaugos reglamentas
PPTF – Pinigų plovimo ir teroristų finansavimas
PPTFPI – Pinigų plovimo ir teroristų finansavimo prevencijos įstatymas
LB – Lietuvos Bankas
FRD – Finansų rinkos dalyvis
FinTech – Financial Technology (liet. Finansų technologijos)
KYC – Know your customer (liet. kliento pažinimo procedūra)
BVP – Bendrasis vidaus produktas
DRF – Distributed Random Forest (liet. Paskirstytas atsitiktinių miškų modelis)
GBM – Gradient Boosting Machine (liet. Gradiento stiprinimo mašina)
LLM – Large language model (liet. Didelis kalbos modelis)
FNTT – Finansinių nusikaltimų tyrimo tarnyba
NLP – Natural Language Processing (liet. Natūralios kalbos apdorojimas)

Bartnykas, Martynas. (2025). A Transaction Monitoring Method Using Process Mining and Large Language Models. MBA Graduation Paper. Kaunas: Kaunas Faculty, Vilnius University. 102 p.

SUMMARY

KEYWORDS: anti-money laundering, transaction monitoring, process mining, large language models, false positive alerts, efficiency improvement, features engineering

Relevance of the topic. The rapid growth of the financial technology (Fintech) sector has significantly increased the volume and complexity of financial transactions, creating new challenges for anti-money laundering and transaction monitoring systems. Financial are required to comply with strict regulatory frameworks while maintaining efficiency and customer-friendly payment processes. Traditional rule-based systems generate a high number of false positive alerts, leading to excessive operational costs, inefficient use of human resources, and reduced overall effectiveness of AML processes. Therefore, the application of advanced analytical methods, such as process mining and large language models, has become increasingly relevant for improving transaction monitoring efficiency.

The object of paper – the efficiency of transaction monitoring process in financial institutions, with a focus on false positive alerts reduction.

Paper's objective – to propose a method that integrates process mining and large language models in order to improve the efficiency on transaction monitoring processes by reducing the number of false positive alerts.

Paper tasks:

- To analyze the principles and limitations of traditional transaction monitoring systems in the context of anti-money laundering.
- To examine the theoretical and practical applications of process mining techniques and large language models in financial crime detection.
- To develop a hybrid method combining process mining and large language models for extracting additional risk-related features from transaction data.
- To conduct an experimental evaluation of the proposed method using transaction monitoring dataset and assess its effectiveness with machine learning models.

Conclusions. The results of the research demonstrate that the integration of process mining and large language models can significantly enhance transaction monitoring systems. Using process mining for identification of complex behavioral patterns and large language models for incorporation of unstructured and rarely used data contributes in generating new features that can more effectively

indicate fraudulent transactions. The proposed hybrid approach contributes to a reduction in false positive alerts and improves the overall efficiency and accuracy of anti-money laundering processes.

Scope of paper. The paper consists of three main parts and comprises 102 pages, including theoretical analysis, method proposal and development, and experimental evaluation. The study includes 11 tables, 14 figures, and 8 appendices. A synthesized transaction data set was used due to data protection constraints, and the experimental analysis was conducted using process mining tools, selected large language model and machine learning models for final evaluation.

IVADAS

Temos aktualumas. Finansinių mokėjimų sektorius sparčiai auga, siūlydamas inovatyvius sprendimus, kurie optimizuoja finansinių paslaugų teikimą ir siekia gerinti vartotojų patirtį. Viena iš svarbiausių finansų įstaigų sričių yra atitikties užtikrinimas, laikantis visų reguliavimo reglamentų ir įstatymų, todėl reguliuojami subjektai privalo užtikrinti aukštą pinigų plovimo prevencijos kokybę. Didžioji dauguma finansų įstaigų susiduria su problema, kuomet inovacijų greitį ir proveržį stabdo griežti įstatymai ar kiti reglamentai, kurių privaloma laikytis. Pavyzdžiui, finansų įstaigos siekia užtikrinti kuo optimalesnį mokėjimų greitį klientams, tačiau mokėjimų stebėsenos sistemos privalo atitikti įstatymuose nurodytus reikalavimus. Dėl šios priežasties mokėjimų stebėsenos sistemos dažnai generuoja daug klaidingai teigiamų įspėjimų, todėl lemia neefektyvų žmoniškųjų išteklių naudojimą, mažina bendrą atitikties vykdymo kokybę, riboja inovatyvumą. Siekiant spręsti šią problemą, būtina taikyti pažangius automatizavimo metodus, tokius kaip finansų proceso gavybą ir didelių kalbos modelių (angl. Large language models, LLMs) panaudojimą. Pritaikant naujus ir inovacijas naudojančius metodus, finansų institucijos galėtų labiau optimizuoti mokėjimų stebėsenos prevencijos procedūrą ir jai naudojamas sistemas, kurios generuotų mažesnę klaidingai teigiamų įspėjimų kiekį bei leistų efektyviau paskirstyti resursus, siekiant užtikrinti bendros atitikties aukštą lygį.

Problemų ištyrimo lygis. Visame finansų sektoriuje mokėjimų stebėsenos sistemos efektyvumas yra vienas iš kertinių atitikties ir verslo aspektų, kuris nulemia mokėjimų įstaigos reguliacinę reputaciją, vartotojų patirtį, verslo rodiklius ir t.t. Tradicinė, taisyklėmis grįsta stebėsenos sistema nors yra paprasta ir patogi, tačiau laikui bėgant nesugeba prisitaikyti prie naujų pinigų plovimo schemų, yra per daug priklausoma nuo žmogaus bei generuoja itin didelius kiekius klaidingų įspėjimų. Todėl, ieškoma pačių optimaliausių būdų didinti mokėjimų stebėsenos efektyvumą, o vienu iš pagrindinių įrankių tam tapo finansų proceso gavybos metodai. Toks metodas leidžia identifikuoti mokėjimų srautuose esančius nukrypimus nuo įprastų finansinės veiklos modelių, o tai yra ypač aktualu siekiant nustatyti sukčiavimo rizikas ir nelegalius veiksmus. Detalesnė ir išsamesnė procesų analizė leidžia efektyviau paskirstyti stebėsenos sistemos išteklius ir padeda koreguoti sistemos logiką siekiant didinti efektyvumą (mažesnis kiekis klaidingų įspėjimų).

Pastaraisiais metais pastebimas vis didesnis didelių kalbos modelių eksperimentinis naudojimas mokėjimų stebėsenos sistemose. Pagrindinė tyrimų sritis apima modelių galimybes konvertuoti finansinius duomenis į nestruktūrizuotą tekstinį formatą, kuris naudojamas anomalijų aptikimui. Pastebima, kad dideli kalbos modeliai leidžia į stebėsenos logiką įtraukti retai naudojamus informacijos laukelius (mokėjimo paskirtys, klientų profilio informacija ir t.t.).

Tokie tyrimai atskleidžia pažangių analitinių metodų potencialą, tačiau praktikoje vis dar trūksta atliktų eksperimentų, kurie apimtų finansų proceso gavybos ir didelių kalbos modelių kombinuotą metodą bei vertintų šio metodo daromą įtaką mokėjimų stebėsenos efektyvumui.

Darbo problema – kaip finansų proceso gavybos metodo ir didelių kalbos modelių integravimas į finansų įstaigų mokėjimų stebėsenos procesus gali padidinti šių procesų efektyvumą, siekiant sumažinant klaidingai teigiamų įspėjimų kiekį?

Darbo objektas – finansų įstaigų pinigų plovimo prevencijos procesų efektyvumo didinimas mažinant klaidingai teigiamų įspėjimų kiekį, taikant finansų proceso gavybos metodą ir didelio kalbos modelio metodo integraciją mokėjimų stebėsenos sistemose.

Darbo tikslas – pasiūlyti metodą, integruojant finansų proceso gavybą ir didelius kalbos modelius, skirtą didinti pinigų plovimo prevencijos procesų efektyvumą, mažinant klaidingai teigiamų įspėjimų kiekį mokėjimų stebėsenos sistemose.

Darbo uždaviniai:

1. Išanalizuoti pinigų plovimo prevencijos stebėsenos sistemų veikimo principus ir identifikuoti pagrindinius klaidingai teigiamų įspėjimų susidarymo veiksnius
2. Išnagrinėti procesų gavybos algoritmų ir didelių kalbos modelių taikymo galimybes mokėjimų stebėsenos kontekste bei apibendrinti jų teorinius ir empirinius taikymo rezultatus sukčiavimo identifikavimui.
3. Pasiūlyti hibridinį procesų gavybos ir didelių kalbos modelių metodą, skirtą finansinių mokėjimų duomenų analizei ir papildomų rizikos savybių (features) išvedimui, siekiant sumažinti klaidingai teigiamų įspėjimų kiekį.
4. Atlikti eksperimentą pagal sudarytą hibridinį metodą, panaudojant paruoštus mokėjimų stebėsenos duomenų rinkinius ir įvertinti jo efektyvumą taikant mašininio mokymosi algoritmą.

Darbo struktūra.

Darbo ir tyrimo metodai. Analizuojant finansų proceso ir didelių kalbos modelių integracijos į mokėjimų stebėsenos procesą teorinį aspektą buvo naudojami šie tyrimo metodai: mokslinės literatūros analizė ir sintezė, dedukcija. Kuriamas pasiūlyto sprendimo prototipas

Darbe naudoti literatūros šaltiniai. Teorinėje darbo dalyje daugiausia naudotasi užsienio bei Lietuvos autorių moksliniais darbais, empiriniais tyrimais susijusiais su pinigų plovimo prevencijos ir mokėjimų stebėsenos efektyvinimu, finansų proceso gavybos taikymu mokėjimų stebėsenoje, didelių kalbos modelių pritaikomumas ir reikšmė mokėjimų stebėsenos efektyvinimui.

Darbo teorinė reikšmė. Darbo teorinė reikšmė pasireiškia tuo, kad darbe atliekamas procesų gavybos ir didelių kalbos modelių taikymo mokėjimų stebėsenai analizė ir apibendrinimas. Atliktas darbas papildo esamą šios srities mokslinę diskusiją, siūlant integruoti skirtingų procesų metu gautų

finansinių duomenų savybes, siekiant didinti mokėjimų stebėsenos efektyvumą, mažinant klaidingai teigiamų įspėjimų kiekį.

Darbo praktinė reikšmė. Darbo praktinė reikšmė grindžiama siūlomu procesų gavybos ir didelių kalbos modelių metodu ir eksperimento metu gautais realiais efektyvumo rezultatais. Gauti naudojamo metodo galutiniai rezultatai (sugeneruotų klaidingai teigiamų įspėjimų kiekis ir kt.) gali būti naudojami finansų institucijose, siekiant optimizuoti mokėjimų stebėsenos procesą, mažinant klaidingai teigiamų įspėjimų kiekį ir didinant analitikų darbo efektyvumą.

Darbo apibojimai ir sunkumai. Darbo atlikimu metu buvo susidurta su sunkumais, ieškant tinkamo finansinių duomenų rinkinio dėl fakto, kad pagal BDAR reglamentą ši informacija yra itin jautri, todėl atliekant darbą buvo naudojamas sintezuotas duomenų rinkinys. Taip pat, naudojant didelį kalbos modelį, buvo susidurta su techninių parametrų ir finansinių kaštų apribojimais, todėl atliekant eksperimentą panaudotas mažiau efektyvus GPT modelis.

Darbo apimtis – darbą sudaro trys dalys, bendrai darbą sudaro 102 puslapiai, iš kurių 74 sudaro minėtas tris darbo dalis. Taip pat darbe pateikiama 11 lentelių, 14 paveikslų ir 8 priedai.

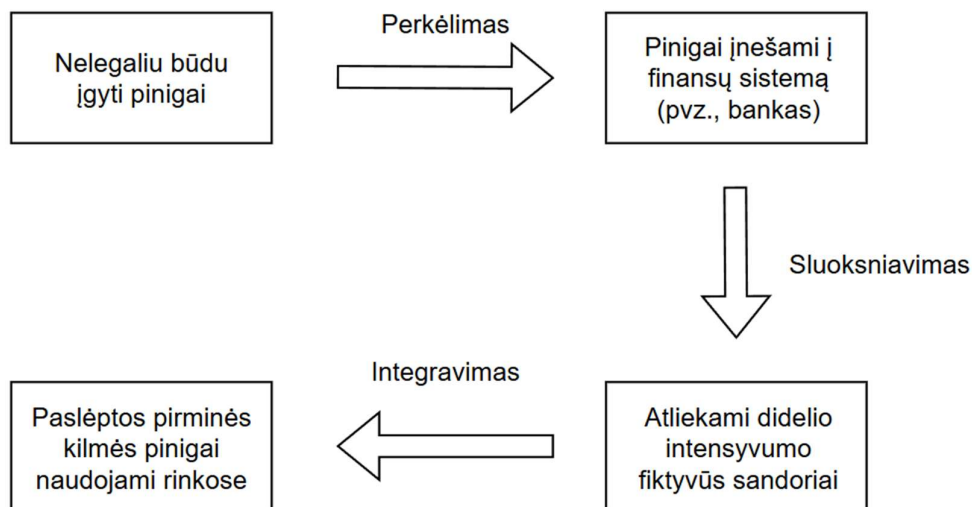
1. PINIGŲ PLOVIMO PREVENCIJOS EFEKTYVUMO DIDINIMO TEORINIAI ASPEKTAI MOKĖJIMŲ STEBĖSENOS KONTEKSTE

Pirmojoje šio darbo dalyje analizuojama pinigų plovimo prevencijos samprata ir mokėjamų stebėsenos veikimo principai, akcentuojant esminius efektyvumo trūkumus. Taip pat analizuojamas finansų proceso gavybos metodo taikymas mokėjamų stebėsenos sistemose, jo taikymo būdai, privalumai ir trūkumai. Galiausiai atliekama didelių kalbos modelių (LLMs) pritaikomumo analizė mokėjamų stebėsenos sistemose, akcentuojant efektyvumo didinimo potencialą (klaidingai teigiamų įspėjimų kontekste).

1.1 Pinigų plovimo prevencijos samprata ir principai

Pinigų plovimu yra laikomas neteisėtas veikimo principas, kuomet nusikalstamu ir nelegaliu būdų įgyjamos lėšos ar turtas yra maskuojamas ir apsimestinai pateikiamas kaip legaliai ir teisėtai įgytas turtas ar pajamos (Gilmour et. al., 2024, p. 8). Tokiu būdu pavieniai nusikaltėliai ar nusikalstamosios grupuotės įgyja fiktyvias teisėtas lėšas, kurios yra naudojamos kitoms nusikalstamoms veikloms finansuoti. Tokiam procesui atlikti dažniausiai naudojamos sofistikuotos ir kompleksiškos schemas, veikiant per trečiąsias šalis, pagrindinius kanalus, naudojantis korumpuotais aukšto rango ir didelę įtaką turinčiais individais ar organizacijomis, todėl prevenciją atliekančios institucijos ne visada sugeba užkirsti kelią „nešvarių“ pinigų ar turto judėjimui. Pinigų plovimo schemų kompleksškumas atsispindi statistikoje – Jungtinių Tautų Narkotikų ir Nusikaltimų prevencijos biuro (angl. United Nations Office on Drugs and Crime) duomenimis, per metus pasaulyje išplaunama tarp 2 ir 5 procento pasaulio BVP, t.y. nuo 800 milijardų iki 2 trilijonų JAV dolerių kasmet (UNODC, 2025). Šie skaičiai atskleidžia ir akcentuoja pinigų plovimo kaip globalios problemos opumą, todėl tokios problemos sprendimui turi būti organizuojamas tarptautinio lygmens koordinuotas atsakas.

Vertinant kasmet išplaunamų pinigų kiekį svarbu akcentuoti, kad tokie pinigų srautai daro didelį poveikį globaliai ekonomikai, todėl tarptautinė finansų bendruomenė ir įvairios nacionalinės priežiūros institucijos kiekvienais metais skiria vis didesnę dėmesį, siekiant sumažinti plaunamų pinigų srautus. Siekiant sumažinti pinigų plovimo mastą tarptautiniu lygiu pasaulyje veikia tokios organizacijos ir institucijos, kaip – Finansinių veiksmų darbo grupė (FATF – angl. Financial Action Task Force), Europos Sąjunga, Jungtinių Tautų organizacija, Užsienio Turto Kontrolės biuras (OFAC – angl. Office of Foreign Assets Control) ir kitos nacionalinės organizacijos (FNIT, Europol ir t.t.). Tam, kad institucijos gebėtų efektyviai kovoti su pinigų plovimu, šis procesas buvo išskaidytas į tris esminius etapus – perkėlimas, sluoksniavimas ir integracija, kurių schema nurodoma 1 paveiksle:



Šaltinis: sudaryta autoriaus

pav. 1 Pinigų plovimo etapų schema

- **Perkėlimas** (angl. Placement) – tai yra pirminis etapas po įvykusios nelegalios veiklos (sukčiavimo, vagystės, kitu neteisėtu būdu įgyti pinigai ar turtas), kuris skirtas atlikti pirminį nešvarių pinigų įvedimą į finansų sistemą. Atliekant šį žingsnį nusikaltėliai įneštas pinigines sumas naudoja perkant itin brangias prekes arba mėgina pinigus išvežti į kitą šalį, o pagrindinis tikslas yra siekis kuo greičiau pakeisti nelegalių pinigų formą, iškeičiant į įvairų turtą, taip siekiant išvengti institucijų dėmesio (Mulig & Smith, 2004). Pinigų perkėlimo etapas tampa pačiu rizikingiausiu, nes toks procesas reikalauja mažomis sumomis pakartotinai įnešti į sąskaitą pinigus (indėlio, pavedimų ar kita forma), todėl gali sukelti įtarimus PPP analitikui, jeigu vykdoma veikla sugeneruos įspėjimą (angl. Alert). Finansinių veiksmų darbo grupė (angl. FATF) pastebi, jog šiame etape naudojami įvairių formų pinigai, pavyzdžiui, nešvarūs grynieji pinigai „įnešami“ per verslus, kurie itin aktyviai naudoja grynuosius (k kazino, nišinės parduotuvės, pinigų perlaidos), tačiau aktyviai naudojami tiek elektroniniai pinigai (sukčiavimo atvejais), tiek intensyvėja kriptovaliutų naudojimas, dėl jų patogumo atsiskaitant nelegaliose svetainėse (FATF, 2018).
- **Sluoksniavimas** (angl. Layering) – etapas atliekamas iš karto po Perkėlimo etapo, kurio metu siekiama kuo labiau įvairiomis, tačiau logiškais ir pagrįstais, piniginių operacijomis paslėpti pirminę lėšų kilmę. Atliekant sluoksniavimo etapą nusikaltėliai veikia didindami atliekamų operacijų skaičius išrašant ar apmokant įvairias sąskaitas-faktūras, atlieka įvairius pavedimus už tam tikrus pirkinius, dalį fiktyvių pinigų „užmaskuoja“ siekiant gauti didesnes

paskolas iš bankų/kredito institucijų, naudoja lėšas didinti fiktyvių įmonių kapitalus ir kt. Taip pat dažnai išnaudojami išvestiniai finansiniai instrumentai, tokie kaip sutartys dėl skirtumo (angl. CFDs – Contract For Difference), perkami įvairūs vertybiniai popieriai, dideliais kiekiais atidaromos investicinės sąskaitos, ypač ofšorinėse jurisdikcijose (Schneider & Windischbauer, 2008). Sluoksniavimui atlikti gali būti naudojamos ir įvairios fiktyvios įmonės, kurios tarpusavyje imituoja ekonominę veiklą, tačiau realybėje pinigai yra pervedami iš vienos sąskaitos į kitą, kuomet galiausiai sukuriamas legaliai įgytų/uždirbtų pinigų įspūdis. Verta paminėti, kad šiame etape jokios paslaugos ar pirkiniai nėra tiesiogiai perkami, o tik atliekamos operacijos, vaidinamos prekių/paslaugų pirkimą, investavimą į išvestinius finansinius instrumentus ir t.t.

- Integravimas (angl. Integration) – paskutinis etapas, kuomet „išplautos“ lėšos yra įliejamos į ekonomiką legaliomis formomis, t.y. atsiskaitant už suteiktas paslaugas ar prekes, įsigyjant įvairios formos kilnojamą ar nekilnojamą turtą ir t.t. Šiame atliekamos operacijos skirtos prabangaus turto ar paslaugų įsigijimui naudojant lėšas su užmaskuotu pirminiu šaltiniu, todėl tokie mokėjimai kelia mažiau klausimų priežiūros institucijomis, t.y. pasiekiamas galutinis tikslas. Verta paminėti, kad trečiame etape nusikaltėlis gali pasinaudoti įvairių patikimų partnerių ar trečiųjų šalių paslaugomis, kurie jo vardu atlieka įvairias pinigines operacijas naudojant išplautus nelegalius pinigus, taip sumažinant riziką pakliūti į teisėsaugos ar institucijų akiratį (FATF, 2018).

Aukščiau paminėta pinigų plovimo etapų schema yra itin bendrinė, kadangi nepaisant esminių etapų, nusikaltėliai formuoja naujas „sub-schemas“, pridedant papildomus žingsnius, įvedant naujus etapus pagal įstatymų ar priežiūros tarnybų spragas, todėl prevencijos tarnybos, tiek nacionaliniu, tiek globaliu lygiu rengia įvairias pinigų plovimo prevencijos gaires ir reikalavimus, kurių institucijos privalo laikytis. Europos Sąjungoje pinigų plovimo prevenciją reglamentuoja ES direktyvos, kurios yra perkeliamos į valstybių narių nacionalinius teisės aktus. Lietuvoje pagrindinis šios srities teisės aktas yra Lietuvos Respublikos pinigų plovimo ir teroristų finansavimo prevencijos įstatymas (LR PPTFPĮ, 2025), kuris nurodo ir įtvirtina aiškius reikalavimus finansų rinkos dalyviams (įskaitant FinTech sektorių) dėl klientų identifikavimo, verslo santykių ir operacijų stebėsenos, rizikos vertinimo ir įtartinų operacijų pranešimo (LR PPTFPĮ, 2025). Tuo tarpu Lietuvos Bankas prižiūri, kad finansų rinkos dalyviai laikytųsi įstatymuose numatytų reikalavimų ir atlieka įvairaus tipo patikrinimus, kurių metu vertinamos šios pinigų plovimo prevencijos įgyvendinimo priemonės:

1. Kliento deramo patikrinimo (KYC – angl. Know Your Customer) vykdymas. Atliekant šią funkciją finansų įstaigos privalo nustatyti kliento naudos gavėjo tapatybę, suprasti dalykinių santykių tikslą, nuolat stebėti kliento veiklą ir atnaujinti informaciją, siekiant užtikrinti, kad sandoriai atitiktų kliento profilį ir rizikos lygį (LBA, 2015). Šis procesas atliekamas tiek

individualiems klientams, tiek juridiniams klientams (ir su juridiniu klientu susijusiems asmenims).

2. Nuolatinė operacijų stebėseną – visi finansų rinkos dalyviai privalo turėti vidaus kontrolės procedūras dalykinių santykių stebėsenai, privalo sustabdyti įtartinas operacijas ir apie jas per numatytą laiko tarpą pranešti FNTT, ypač jei sandoriai yra neįprasti, sudėtingi ar susiję su aukštos rizikos šalimis (LB, 2022). Verta paminėti, kad operacijų stebėseną būtina atlikti, siekiant ar tinkamai buvo identifikuota kliento veikla ir vykdoma veikla po verslo santykių užmezgimo fakto.
3. Rizikos pagrįstumo (vertinimo) principas – tiek kliento deramo pažinimo metu, tiek nuolatinės operacijų stebėsenos metu būtina užtikrinti atitinkamą rizikos vertinimo matricą išskiriant įvairių sričių keliamas rizikas (geografinė rizika, verslo rizika, produkto rizika ir kt.). Tokia priemonė būtina siekiant tinkamai suvaldyti rizikas, kurios kyla aptarnaujant tam tikrus subjektus, kurie vykdo veiklą itin rizikinguose sektoriuose ar itin rizikingose šalyse. Atitinkamai pagal identifikuotą rizikos lygį finansų rinkos dalyvis privalo taikyti sustiprintą deramą kliento pažinimą ir kitus papildomus procesus pagal Lietuvos Banko siūlomus gerosios praktikos motyvus.
4. Informacijos saugojimas ir įtartinų operacijų pranešimas – kaip minėta aukščiau, finansų rinkos dalyvis vykdydamas nuolatinę operacijų stebėseną ir identifikavęs įtartiną veiklą, privalo atitinkamai informuoti FNTT. Taip pat visi įpareigoti subjektai privalo tvarkyti įvairius registracijos žurnalus, susijusius su klientų operacijomis, įtartinais sandoriais, dalykiniais santykiais ir jų nutraukimu, bei saugoti šiuos duomenis nuo 5 iki 8 metų, priklausomai nuo informacijos pobūdžio. Finansų rinkos dalyvių reikalaujama išsaugoti klientų tapatybę patvirtinančius dokumentus, korespondenciją ir sandorius įrodančius duomenis (PPTFPĮ, 2025).
5. Vidaus kontrolės mechanizmai – tokie mechanizmai finansų rinkos dalyviams yra būtini, siekiant užtikrinti, kad juridiniai asmenys laikytųsi teisės aktų, tinkamai saugotų turtą nuo sukčiavimo ar išvaistymo, veiklą vykdytų laikydamasis finansų valdymo principo. Taip pat FRD privalo turėti atskiras vidaus kontrolės procedūras (ir scenarijus), kad laiku pastebėtų teroristų finansavimo atvejus ir pinigų plovimo atvejus (Lietuvos Bankas, 2022).

Aukščiau paminėti yra pagrindinės pinigų plovimo prevencijos priemonės, skirtos užkirsti kelią galimiems pinigų plovimo ir teroristų finansavimo atvejams, o šių priemonių įgyvendinimą prižiūri atitinkamos nacionalinės institucijos ar reguliatorius.

Taigi, apibendrinant, poskyryje buvo atlikta pinigų plovimo reiškinių analizė, apibrėžta jo samprata, įvardintas tarptautinis ekonominis mastas bei išskirti pagrindiniai veikimo etapai – perkėlimas, sluoksniavimas ir integracija. Taip pat detalai išskirtos pagrindinės prevencinės

priemonės, kurias taiko finansų rinkos dalyviai Lietuvoje, vadovaudamiesi teisės aktais ir priežiūros institucijų gairėmis. Esminis dėmesys skirtas kliento pažinimo principams, sandorių stebėsenai, rizikos vertinimui, informacijos saugojimui bei vidaus kontrolės mechanizmų svarbai. Aptarti tiek nacionalinio, tiek tarptautinio reguliavimo aspektai leidžia geriau suprasti, kokio sisteminio požiūrio reikalauja pinigų plovimo prevencijai ir kokie praktiniai veiksmai taikomi šioms rizikoms suvaldyti finansų sektoriuje.

1.2 Mokėjimų stebėsenos principai ir efektyvumo problematika

Aukščiau minėtame skyriuje nuolatinę mokėjimų stebėseną įvardijome kaip vieną iš kertinių veiksmų, kuris yra itin svarbus siekiant užkirsti pinigų plovimo veiklą, sutrikdyti galimas plovimo schemas, sukelti nepatogumus ir galiausiai – identifikuoti galimus pinigų plovimo subjektus. Mokėjimų stebėsenos svarbą ir reikšmingumą pabrėžia tiek Europos Sąjungos Pinigų Plovimo Prevencijos direktyvos, tiek šalių įstatymai nacionalinių lygmeniu, pavyzdžiui, Pinigų Plovimo ir Teroristų Finansavimo Prevencijos Įstatymas Lietuvoje. LB atliktoje apžvalgoje indikuoja, kad finansų rinkos dalyviai (FRD) privalo vykdyti ir įgyvendinti nuolatinę klientų dalykinių santykių ir operacijų stebėseną, siekdami užtikrinti, kad kliento veikla ir atliekamos operacijos atitiktų finansų institucijos turimas žinias apie kliento verslą, sudarytą rizikos profilį ir lėšų kilmę (Lietuvos Bankas, 2022). Taigi, dėl šių priežasčių FRD privalo turėti tinkamus techninius pajėgumus ir infrastruktūrą, kuri leistų efektyviai vykdyti kliento ir jo operacijų stebėseną, o tam gali būti naudojamos tiek vidiniai sprendimai, tiek iš trečiųjų tiekėjų perkant ir integruojant techninius sprendimus.

Nepaisant diegiamų sistemų ar kitų mokėjimų stebėsenos techninių produktų, verta paminėti, kad atliekamas procesas, anot Lietuvos Banko, turi remtis šiais šešiais kartiniais aspektais:

1. Momentinė stebėseną – šio tipo stebėseną vykdoma taip, kad aukštos rizikos operacijos būtų realiu laiku automatiškai neįvykdytos, kol atsakingas darbuotojas neįvertina sugeneruoto įspėjimo ir pačios operacijos teisėtumo. Dažniausiai taisyklėmis paremta sistema atlieka loginį vertinimą pagal iškeltas sąlygas, siekiant užkirsti ir maksimaliai sumažinti nelegalių lėšų judėjimą tarp finansinių institucijų. Taip pat verta pabrėžti, kad pagal 6-tąją AML direktyvą, įmonės privalo taikyti sustiprintą kliento patikrinimą prieš mokėjimo paleidimą, itin aukštos rizikos klientams, taip mažinant riziką susijusią su nelegaliu lėšų judėjimu (European Parliament & Council, 2018).
2. Retrospektyvi stebėseną – tuo tarpu retrospektyvi stebėseną nestabdo realiu laiku kliento vykdomų mokėjimų, tačiau atlieka jau įvykdytų kliento operacijų peržiūrą per iš anksto nustatytą laikotarpį (pvz., paskutinį mėnesį ar ketvirtį). Šios procedūros tikslas yra patikrinti kliento operacijų atitikimą pagal KYC (angl. Know Your Customer) profilio informaciją,

ekonominės veiklos sektorių ir deklaruotus finansinius srautus. Papildomai, naujausia AML direktyva įpareigoja įmones retrospektyvios stebėsenos metu identifikuotus įtartinus mokėjimus pagal tvarką raportuoti atitinkamoms šalies institucijoms (European Parliament & Council, 2018), pavyzdžiui, Lietuvos atveju – Finansinių Nusikaltimų Tyrimų Tarnyba (FNTT).

3. Sandorių tipų ir klientų segmentavimo kriterijai – siekiant optimizuoti stebėsenos našumą, FRD savo klientus skirsto pagal PPTF rizikos lygį, geografinę lokaciją, ekonominę veiklą ir sąskaitų tipus. Turint tokią informaciją, galima geriau pritaikyti stebėsenos scenarijus, kuriais siekiama nustatyti klientui neįprastą ar įtartiną veiklą (Lietuvos Bankas, 2022), tačiau verta pastebėti, jog stebėsenos sistemos optimaliam veikimui FRD turėtų pasirinkti kliento profilio kokybę ir kiekybę, kad scenarijams būtų galima pritaikyti ir lyginamąją analizę.
4. Stebėsenos apimtis ir įrankiai – įprastai mokėjimų stebėseną apima įeinančių ir išeinančių operacijų srautus, tačiau LB pastebi, kad FRD privalo stebėti ir neaktyvias (angl. dormant) ir tranzitines sąskaitas, kurios dažnai yra naudojamos pinigų plovimo etapuose, o be šių aspektų reguliatorius skatina atlikti gilesnes ir nuodugnesnes analizes apie klientų veiklą, naudojant sąsajų ieškojimo (angl. Networks and links) metodą ir pan. (Lietuvos Bankas, 2022). AML direktyva pabrėžia, kad atliekant mokėjimų stebėseną FRD papildomai privalo turėti prieigą prie centralizuotų registrų ir kitų svarbių duomenų bazių, kurios palengvintų mokėjimų stebėsenos proceso vykdymą, leistų patvirtinti ar paneigti kliento pateiktą informaciją naudojantis patikimais ir nepriklausomais šaltiniais (European Parliament & Council, 2018).
5. Rizika grindžiamas metodas (angl. Risk-based approach) – atliekamos stebėsenos intensyvumas ir dažnis turi tiesiogiai koreliuoti su kliento rizikos įverčiais, todėl FRD privalo turėti išsamius klientų KYC profilius, juos tęstinai atnaujinti ir perskaičiuoti rizikos įverčius bei pagal juos koreguoti mokėjimų stebėsenos intensyvumą. Taip pat FRD turi nuolat reaguoti į kintančius PPTF rizikos veiksnius, adaptuotis atsižvelgiant į naujus reguliacinius reikalavimus, tarptautines rekomendacijas ir atliekamus auditus. Lietuvos Bankas (2022) pažymi, kad FRD privalo reguliariai atlikti imčių testavimą (angl. Sample back testing), kad įvertinti sistemos tikslumą ir koreguoti parametrus. Toks reikalavimas užtikrina, kad mokėjimų stebėsenos sistema išlieka jautri naujoms grėsmėms.
6. Vėlavimų ir kokybės kontrolė – Lietuvos Banko pateiktoje analizės ataskaitoje pabrėžiama, kad visi automatiškai sugeneruoti signalai turi būti peržiūrėti per nustatytą terminą, o vėluojančių atvejų stebėseną turi būti fiksuojama ir analizuojama, siekiant identifikuoti procesų tobulinimo galimybes (Lietuvos Bankas, 2022). Tokių terminų įgalinimas skatina finansų rinkos dalyvius efektyviau valdyti išteklius ir turėti vidinius procesus, kad būtų laikomasi reguliatoriaus nustatytų reikalavimų.

Visi šie aspektai yra kartiniai norint tinkamai ir kokybiškai atlikti mokėjimų stebėseną, todėl itin svarbu, kad finansų rinkos dalyvis ne tik vykdytų visus keliamus reikalavimus, bet ir skirtų pakankamą finansavimą procesų įgyvendinimui bei skirtų pakankamai žmogiškųjų išteklių.

Nepaisant aktyvaus finansų rinkos reguliavimo ir finansų rinkų dalyvių priežiūros bei nuolatinio direktyvų ir įstatymų tobulinimo, pinigų plovimo prevencijos sistemos, tame tarpe ir mokėjimų stebėseną susiduria su technologiniais efektyvumo iššūkiais. Vienas iš pagrindinių su mokėjimų stebėseną susijusių efektyvumo spragų/problemų yra klaidingi įspėjimai (angl. False positive alerts), kurie yra sugeneruojami visiškai legaliems ir teisėtiems mokėjimams. Vertinant dabartinius atliktus tyrimus ir kitą akademinę literatūrą, galime teigti, kad klaidingai sugeneruoti įspėjimai sudaro tarp 90 proc. ir 95 proc. visų sugeneruojamų įspėjimų (Oztas et. al., 2024). I. Vorobyev & A. Krivitskaya (2022) taip pat pastebi, kad žemo tikslumo arba aukšto klaidingų įspėjimų rodiklio problema geriausiai atspindi statistikoje, kuri rodo, kad vidutiniškai iš penkių užblokuotų mokėjimų tik vienas pasitvirtina kaip sukčiavimo atvejis, o kas šeštas vartotojas yra klaidingai užblokuojamas dėl sistemų neefektyvumo ir klaidingų išvadų. Viena iš esminių priežasčių, lemiančių mokėjimų stebėsenos technologinį neefektyvumą yra taisyklėmis grindžiama sistema, kuri sukuria įspėjimą pagal nustatytas taisyklių sąlygas. Dažnai finansų įstaigos naudoja taisyklėmis grindžiamas sistemas dėl jų skaidrumo ir paprastumo – taisyklės yra aiškios, patikrinamos, lengvai suprantamos, todėl finansų įstaigos gali realiu laiku aptikti įtartinas veiklas, o audito metu nesunku įrodyti, kad finansų įstaiga laikosi visų reguliavimo reikalavimų (Omoseebi, 2025). Tačiau, verta paminėti, kad sistemos paprastumas sukelia gana didelius efektyvumo kaštus – sistemą specialistai turi atnaujinti rankiniu būdu, ji negali prisitaikyti prie kintančių mokėjimų tipologijos, naujos taisyklės turi būti pasiūlomos pačio specialisto, binarinė sąlygų logika neleidžia įsigilinti į mokėjimų kontekstą. Ketenci et. al. (2021) pažymi ne tik aukštą (virš 90 proc.) generuojamų klaidingai teigiamų įspėjimų problematiką, bet ir faktą, kad taisyklėmis grindžiama sistema yra visiškai priklausoma nuo žmogaus, todėl jos veiksmingumas priklauso nuo specialisto patirties, o dėl žmogiškojo faktoriaus kyla grėsmė, kad taisyklių rinkiniai bus atskleisti nusikaltėliams.

Analizuojant literatūrą, galima pastebėti, kad klaidingų įspėjimų kiekis yra itin opi ir vis dar sprendžiama problema, kuri atspindi proporcingai lyginant klaidingus įspėjimus ir visus sugeneruotus įspėjimus. Taip pat pabrėžiama, kad nepaisant fakto, jog taisyklėmis grindžiama mokėjimų stebėseną yra vienas iš paprasčiausių ir populiariausių sprendimų vykdant mokėjimų stebėseną, tačiau būtent ši sistema kelia didžiausią iššūkį mažinant klaidingų įspėjimų kiekį, dėl savo neefektyvumo, priklausomybės nuo specialistų ir limituotos binarinės logikos. Galiausiai, paprasta loginė sistemos struktūra sunkiai priderinama prie besikeičiančių klientų elgsenos modelių.

Finansų technologijų rinkos ir jos siūlomų produktų populiarumas itin stipriai išaugo per pastaruosius dešimt metų – tai lemia žymiai didesnius pinigų perlaidų srautus tarp elektroninių pinigų

institucijų ir kitų tradicinių finansų rinkos dalyvių (bankai, kredito unijos, t.t.), didėja mokėjimų dažnis, o su juo didėja ir klaidingų įspėjimų problematika. Dėl didelio klaidingų įspėjimo kiekio finansų institucija gali neturėti pakankamai žmogiškųjų išteklių, kurie susitvarkytų su kylančiais iššūkiais ir užtikrintų sistemos greitį, patikimumą ir kliento gerą patirtį. Didelis klaidingų įspėjimų kiekis verčia darbuotojus dirbti stresinėje aplinkoje, todėl specialistai, susidūrę su didelėmis darbo apimtimis, laikui bėgant įpranta atsainiau žiūrėti į sistemos įspėjimus ir tampa mažiau jautrūs tikriems įtartiniams atvejams (Ketenci et. al., 2021). Dėl didelių kiekių klaidingų įspėjimų finansų įstaigos darbuotojai negali užtikrinti kokybiškos ir efektyvios analizės, to pasekoje tampa sunkiau komunikuoti su šalies valstybinėmis institucijomis, kurios yra atsakingos už finansinių nusikaltimų tyrimus. Šiai problemai pašalinti finansų įstaiga susiduria su papildomomis išlaidomis – Oztas et. al. (2024) pabrėžia, kad klaidingi įspėjimai išaugina atliekamų tyrimų kiekį, kuris išaugina kaštus, todėl didėja spaudimas keisti dabartinius stebėsenos metodus, tačiau tai taip pat kelia papildomus didelius kaštus. Statistiškai, atliekamas tyrimas reikalauja sąveikos su klientu (skambučio ar laiško forma), o tai, remiantis Europos Centrinio Banko duomenimis, sukuria papildomus kaštus nuo 1,5 euro iki 5 eurų už kiekvieną užblokuotą mokėjimą (Vorobyev & A. Krivitskaya, 2022).

Tolimesnė literatūros analizė parodė, kad klaidingų įspėjimų problema kyla ir dėl prasto kliento KYC duomenų tvarkymo, finansų įstaigos departamentų, atsakingų už pinigų plovimo prevenciją, izoliuotumas ir skirtingi reguliaciniai režimai tarpvalstybiniu lygiu. Oztas et. al. (2023) pastebi, kad prasta duomenų kokybė taip pat gali lemti didesnius klaidingai generuojamų įspėjimų kiekius – netikslūs ir neišsamūs duomenys gali trukdyti veiksmingai taikyti sandorių stebėjimo metodus. Tikslūs ir išsamūs KYC duomenys apie klientą gali ne tik padėti greičiau išspręsti ir užbaigti klaidingo įspėjimo tyrimą, bet ir išvengti tyrimo pradėjimo. Autorius taip pat pabrėžia, kad iššūkiai gali kilti ir dėl departamentų, užtikrinančių pinigų plovimo prevenciją, izoliuotumo ir savarankiškumo atliekant įvairias užduotis, ko pasekoje KYC metu surinktą informaciją prastai ar netiksliai įsisavina kiti departamentai (Oztas et. al., 2024). Praktikoje ši problema atsispindi vertinant skirtingų AML srities departamentų pareigybių aprašus, kurie neretai tik teoriškai persidengia tam tikromis pareigybėmis, tačiau praktikoje atlieka visiškai skirtingas funkcijas, todėl prastėja informacijos pasidalijimas, krenta komunikacijos lygis. Galiausiai, verta paminėti, kad reguliaciniai reikalavimai finansų institucijoms nuolat keičiasi ir globaliu lygmeniu matomas ryškus disbalansas, ko pasekoje trūksta aiškaus tarptautinio standarto, kuris nustatytų aiškias mokėjimų stebėsenos taisykles.

Taigi, apibendrinus galime teigti, kad nuolatinė mokėjimų stebėseną yra esminė priemonė, leidžianti ne tik užkirsti kelią pinigų plovimo schemoms, bet ir laiku identifikuoti galimus pažeidimus, derinant techninius sprendimus su nuoseklia stebėsenos logika. Siekiant užtikrinti stebėsenos efektyvumą, FRD privalo turėti lanksčią infrastruktūrą, kuri gebėtų realiuoju laiku blokuoti aukštos

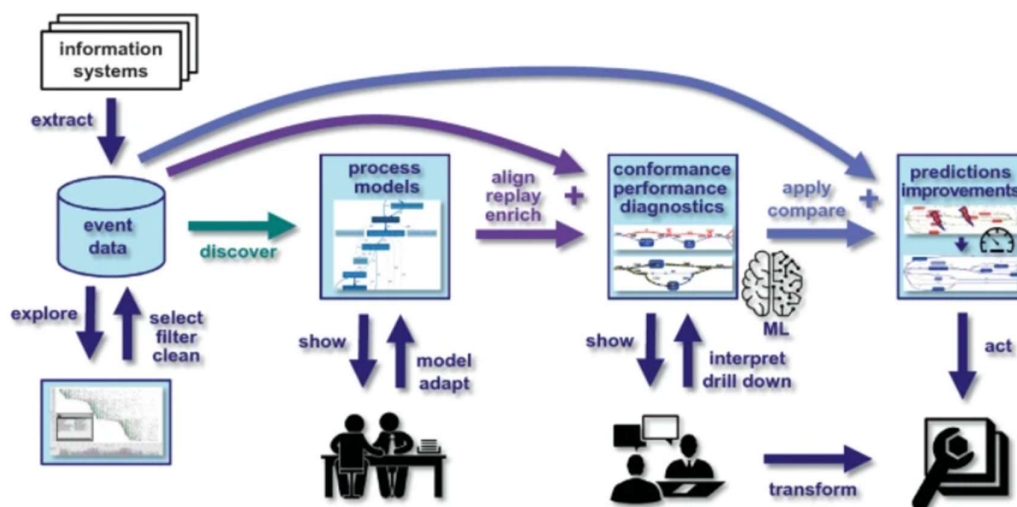
rizikos operacijas, atlikti retrospektyvią analizę ir pritaikyti segmentinius stebėjimo scenarijus pagal kliento rizikos profilį bei stebėti ir neaktyvias ir tranzitines sąskaitas. Rizika pagrįstas požiūris, nuolatinis KYC duomenų atnaujinimas ir sistemos parametrų koregavimas užtikrina, kad kontrolės mechanizmai išliktų jautrūs ir prisitaikę naujoms grėsmėms. Atsižvelgiant į šiuos reikalavimus, matome, kad akademinėje literatūroje vyrauja panašus požiūris, kad taisyklėmis grįsta stebėsenos sistema nebeatitinka keliamų lūkesčių ir rinkos standartų, o generuojamų klaidingų išpėjimų kiekis ženkliai didina institucijos žmogiškųjų ir finansinių išteklių kaštus. Dėl priklausomybės nuo žmogiškojo faktoriaus ir negebėjimo prisitaikyti prie naujų pinigų plovimo schemų, taisyklėmis grįstą sistemą privalo pakeisti pažangesni metodai, kurie stiprintų stebėsenos infrastruktūrą, taupyty įstaigos išteklius ir gerintų efektyvumo rodiklius.

1.3 Finansų proceso gavybos taikymas mokėjimų stebėsenos sistemose

Finansų proceso gavybos metodai ne kartą buvo pritaikyti mokėjimų stebėsenos efektyvumo didinimo kontekste, o atliekami eksperimentai ir toliau analizuojami moksliniai straipsniai parodo, kad tokių metodų taikymas gali daryti teigiamą įtaką sprendžiant mokėjimų stebėsenos efektyvumo trūkumus. Pavyzdžiui, proceso gavyba gali padėti efektyviau panaudoti žmogiškuosius išteklius, nustatyti prioritetus vykdant mokėjimų stebėseną, pagerinti klaidingai teigiamų išpėjimų kiekio rodiklį, sumažinti finansines išlaidas (Ye, 2007).

Procesų gavyba gali būti apibūdinama kaip procesas, kuriuo siekiama pagerinti operacinius procesus sistemingai naudojant įvykių duomenis, kuriuos kombinuojant su procesų modeliu galima gauti vertingų proceso įžvalgų, identifikuoti neatitikimus, nukrypimus, parodo veiklos ir atitikties problemas, pasikartojančius procesus (van der Aalst, 2022). Procesų gavybos principas gali būti orientuojamas į ateitį, pavyzdžiui, prognozuoti vykdomo proceso apdorojimo laiką bei gali būti taikomas praeities procesų analizei – nustatyti procesų trūkumus, silpnąsias verslo ar kitų procesų grandies vietas. Šiam procesui atlikti reikalinga detali ir nenutrūkstama įvykių seka (įvykių žurnalai, laiko žymos ir t.t.), todėl šios technikos taikymas gali būti pritaikomas ir mokėjimų stebėsenos srityje – stebėsenos proceso analizė ir nukrypimai gali padėti nustatyti finansinės elgsenos modelius, generuoti naujas savybes, kurios leistų teisingai nustatyti didesnį kiekį sukčiavimo/pinigų plovimo atvejų. Tačiau R. Accorsi ir J. Lebherz (2022) pastebi, kad užtikrinti efektyvų procesų gavybos taikymą yra būtina tinkamai paruošti duomenų rinkinius ir aiškiai suformuluoti tikslus (konkretūs procesai ar rodikliai) bei išvardinti reikšmingus rodiklius. Siekiant taisyklingai pritaikyti procesų gavybos metodą identifikuojant sukčiavimo mokėjimų struktūras bei atrasti sąsajas, kurias praleidžia įprastos, taisyklėmis grįstos sistemos, yra būtina suprasti šio proceso pagrindinius žingsnius ir taikomus algoritmus.

Klasikinis procesų gavybos pavyzdys, pavaizduotas 2 paveiksle, apima tris pagrindinius žingsnius – procesų atradimą (angl. Discovery), atitikties tikrinimą (angl. Conformance checking) bei praturtinimą (angl. Enhancement).



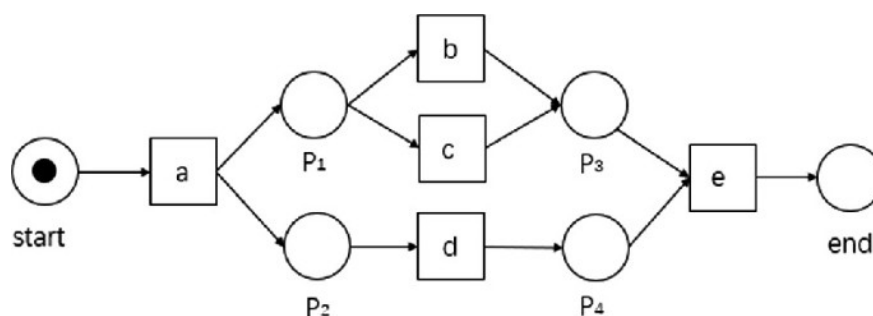
Šaltinis: Aalst, W. M. P. (2022). *Process Mining: A 360 Degree Overview*, p. 5

pav. 2 Procesų gavybos etapų schema

Procesų atradimo etapas įprastai yra pats pirmas žingsnis (atlikus įvykių duomenų rinkinio paruošimą), kurio metu algoritmai automatiškai sukuria proceso modelį, apibūdina stebimą elgesį, o sudarytas modelis yra naudojimas įvairių nuokrypių ir proceso „butelio kakliukų“ nustatymui (W. M. P. van der Aalst, 2022). Tuo tarpu atitikties tikrinimo metodai suteikia mechanizmus, kurie leidžia susieti modeliuojamą ir stebimą realų elgesį, todėl galima nustatyti neatitikimus tarp procesų vykdymo ir procesų modelių, kurie formalizuoja tikėtiną elgesį (J. Carmona, B. Van Dongen, M. Weidlich, 2022). Toks proceso funkcionalumas gali būti itin naudingas vykdant mokėjimų duomenų analizę, kuomet žinomas aiškus atitikties procesas, kokie mokėjimų tipai turėtų būti stabdomi ir analizuojami pagal atitikties ir įstatyminio reguliavimo bazę bei keliamus reikalavimus. Galiausiai atliktos analizės metu surinkti duomenys yra naudojami esamo proceso detalių korekcijai, detalizuojant išteklių, laiko, proceso eigos perdėliojimą ir papildymą. M. de Leoni (2022) teigia, kad akademinė literatūra išskiria du procesų tobulinimo tipus: proceso išplėtimas, kuris yra orientuotas į pagrindinę savybę (angl. Feature), siekia įtraukti skirtingas perspektyvas bei dažniausiai apibrėžia tik kontrolės srauto perspektyvą. Tuo tarpu antras tipas yra procesų tobulinimas, kuris siūlo didžiausią dėmesį suteikti kitoms modelio savybėms ir užtikrinti, kad modelis labiau atspindėtų realybę papildant esamą modelį minėtomis kitomis savybėmis arba apribojant modelį, kad būtų atliekami tik tikėtini procesai. Toks metodo taikymas mokėjimų stebėsenos sistemose galėtų leisti efektyviau

peržiūrėti „teorinės“ prevencinės logikos veikimą realių mokėjimų kontekste ir leistų greičiau reaguoti į nustatytus trūkumus.

Atliekant procesų gavybos eigą pagal aukščiau paminėtus žingsnius, būtina suprasti šių procesų metu naudojamų algoritmų veikimo principus ir jų technologinių savybių pritaikymo galimybes praktikoje. Analizuojama literatūra dažniausiai išskiria šiuos naudojamus algoritmus – Alfa (angl. α -algorithm), Euristinis (angl. Heuristic miner), Induktyvus (angl. Inductive miner) ir Regioninis (region-based). Alfa algoritmas yra laikomas vienas iš pirmųjų procesų gavybos metodų, kurio paskirtis yra rekonstruoti priežastingumo ryšius tarp veiklų, remiantis įvykių sekų duomenimis. Algoritmas analizuoja įvykių žurnaluose (angl. Event logs) registruotų įvykių eiliškumą ir nustato, kurie įvykiai paprastai viena po kitos, o kurios gali vykti ir lygiagrečiai, t.y. atkuria priežastingumą iš įvykių sekų rinkinio (Weerapong et al., 2012). Algoritmas, remiantis aptiktais priežastiniais ryšiais konstruoja Petri tinklą, kurio pavyzdys yra pateikiamas 3 Paveiksle, pagal Petri tinklų klases, kurios reprezentuoja struktūrizuotą darbo ar proceso eigos tinklą (Carmona, 2010).



Šaltinis: Yin, F., Cao, J., Chu, J., ir Olga, S. (2021). *Warehousing Process Mining Research Based on Petri net*, p. 56.

pav. 3 Petri tinklo pavyzdys

Tuo tarpu Euristinis (angl. Heuristic miner) algoritmas Alfa algoritmo atžvilgiu yra pažangesnis ir leidžia efektyviau dirbti su triukšmingais duomenimis. Šis algoritmas kuria proceso modelį analizuodamas kaip dažnai veiklos eina viena po kitos realiuose įvykių žurnaluose bei matuoja šių ryšių stiprumą naudodamas priklausomybės vienetą (angl. Dependency Measure) (Weber, Bordbar ir Tino, 2013). Tuomet algoritmas sujungia dažnai kartu pasikartojančias veiklas, kad sudarytų priklausomybės grafiką (angl. Dependency Graph), o galiausiai identifikuoja, ar proceso šakos atspindi pasirinkimus ar lygiagrečias trajektorijas bei filtruoja triukšmą naudojant slenkstinius nustatymus, kad būtų įtraukti tik reikšmingiausi modeliai (Weber et. al., 2013). Toks veikimo principas leidžia tiksliau nustatyti ryšius tarp įvykių ir procesų pasikartojimą, todėl yra geriau pritaikomas intensyvių ir kompleksinių žurnalų įvykiuose, pavyzdžiui analizuojant mokėjimų elgseną ir ieškant neatitikimų, anomalijų ir įtartinių mokėjimų.

Kitas gana dažnai procesų gavyboje naudojamas algoritmas yra indukcinis (angl. Inductive Miner), naudojamas atrasti procesų modelius iš įvykių žurnalų. Pagrindinė technikos esmė yra paremta įvykių žurnalų skaidymu į mažesnius įvykių žurnalus, kuriant bendrą viso proceso modelį/grafiką. Algoritmas rekursyviai konstruoja proceso medį ir kiekvieno rekursyvo metu identifikuoja „pjūvį“ ir taip dalija įvykių žurnalą iki taško, kuomet pasiekiamas bazinis atvejis (angl. Base case), o jeigu „pjūvis“ nerandamas, tuomet vykdomas duomenų padalijimas (angl. Split), kad būtų galima tęsti modelio kūrimą (Detten, Schumacher ir Leemans, 2023). Kuomet algoritmas negali rasti „pjūvio“ ir atliekamas duomenų padalijimas, tai dažniausiai reiškia, kad analizuojamas įvykių žurnalas yra per daug kompleksiškas, tačiau šiuo atveju išryškėja algoritmo pranašumas prieš kitus algoritmus dėl galimybės technologiškai apdoroti pateiktą duomenų rinkinį. Toks algoritmo veikimo principas leidžia automatiškai aptikti dažnai pasikartojančių procesų ryšius, identifiкуoti nestandartinius veiksmus ir parodyti dažnai pasikartojančius įvykių dubliavimus, kas mokėjimų kontekste galėtų „sufleruoti“ sukčiavimo požymius.

Vis dėlto, mokėjimo srautų duomenys tampa vis kompleksiškesni dėl didėjančių finansinių produktų pasiūlos rinkose, o su šia tendencija sukčiavimo metodai prisitaiko ir išnaudoja šią susidariusią terpę pinigų plovimui vykdyti, todėl tokiose situacijose indukcinis metodas gali generuoti per sunkiai interpretuojamus modelius. Itin sudėtingiems atvejams tampa patrauklesnis neapibrėžtas algoritmas (angl. Fuzzy miner), kuris leidžia supaprastinti ir vizualiai apibendrinti procesų elgseną, išryškinant svarbiausias veiklų sekas ir praleidžiant mažiau reikšmingas. Realiam pasaulyje procesai yra mažiau struktūrizuoti nei tikimasi, todėl įprastiniams kasybos metodams nestruktūrizuoti duomenys kelia rimtų iššūkių formuojant modelius. Tuo tarpu neapibrėžtas algoritmas yra konfigūruojamas ir leidžia gauti skirtingas ir supaprastintas tam tikro proceso dalis, o tam atlikti naudojamas žemėlapių (angl. Roadmap) konceptas (Gupta, 2014). Anot Gupta (2014), atliekant įprastą procesų gavybos procesą gaunami modeliai yra sunkiai suprantami dėl spagečių tipo (angl. Spaghetti-like) problematikos, todėl neapibrėžtas algoritmas šiuo atveju yra tinkamas mažiau struktūrizuotiems procesams, nes taikomos įvairios technikos, pavyzdžiui, ribinių reikšmių šalinimas, susijusių mazgų grupavimą į vieną bendrą mazgą, izoliuotų mazgų naikinimas, leidžia sukurti glaudesnius ir suprantamesnius modelius. Remiantis išskirtomis savybėmis, neapibrėžtas modelis gali būti pritaikomas itin kompleksiose mokėjimų ekosistemose, kurios dažnai generuoja itin didelius kiekius finansinių įvykių, o algoritmo gebėjimas filtruoti triukšmą ir kitus nereikšmingus įvykius gali padėti efektyviau ir tiksliau nustatyti sukčiavimo elgseną ir tendencijas.

Apibendrinant galime teigti, kad procesų gavybos metodai gali būti taikomi mokėjimų stebėsenos ir sukčiavimo prevencijos kontekste, siekiant rekonstruoti ir suprasti realių finansinių procesų eigą bei identifiкуoti reikšmingus nukrypimus, kurie galėtų padėti identifiкуoti rizikingą ar neteisėtą veiklą. Aptariamose procesų gavybos algoritmų techninės savybės leidžia suprasti jų

skirtingą detalumą, tikslumą ir lankstumo lygį, todėl algoritmo pasirinkimas turi priklausyti nuo turimų duomenų struktūros, sudėtingumo ir pan. Alfa algoritmas pasižymi aiškiu priežastingumo ryšių nustatymu, Euristinis algoritmas geba apdoroti triukšmingus įvykių žurnalus bei išskirti reikšmingiausius elgsenos modelius. Tuo tarpu Indukcinis algoritmas sistemiškai skaido sudėtingus procesus į logiškai susietas dalis, o Neapibrėžtas algoritmas geba vizualiai pateikti paprastai suprantamą modelį ir efektyviau filtruoja informaciją, kuomet procesai yra itin kompleksiški.

Praktikoje, mokėjimų stebėsenoje ir kitoje su finansais susijusiose prevencinėse srityse (auditai, viešieji pirkimai, vidinių piniginių srautų valdymas ir t.t.) dažniausiai taikomas neapibrėžtas (angl. Fuzzy miner) algoritmas dėl savo filtravimo savybių ir galimybės pateikti aiškiai suprantamą modelį žmogaus analizei. Baader ir Krcmar (2018) siūlomas klaidingai teigiamų įspėjimų mažinimo prototipas yra paremtas Celonis procesų gavybos įrankiu, kuris naudoja neapibrėžtą algoritmą. Šis pasirinkimas grindžiamas faktu, kad mokėjimų įvykiai generuoja didelius žurnalus, kuriuose yra mažai lygiagretumo, todėl būtent šis algoritmas tampa labiausiai tinkamu. Pateikiamas prototipas paremtas raudonų vėliavų (angl. Red flags) ir procesų gavybos kombinacija nepaisant gana žemo teisingai identifikuoatų sukčiavimo atvejų (48.38%), parodė itin mažą klaidingai teigiamų įspėjimų rodiklį - vos 0.37% (Baader & Krcmar, 2018). Tuo tarpu autorių (Sarno, Dewandono, Ahmad, Naufal ir Sinaga, 2015) pritaikytas neapibrėžto algoritmo ir asocijuotų taisyklių mokymosi kombinuotas metodas parodė bendrą metodo tikslumą 86.5%, aptinkant sukčiavimo atvejus mokėjimuose. Atlikti tyrimai parodo, kad procesų gavybos algoritmai, o ypač neapibrėžtas algoritmas, puikiai dera su kitais duomenų analizės metodais, veikiantis kaip pirmas žingsnis analizuojant ypač didelius ir kompleksinius įvykių žurnalus. Tam antrina ir Jans, Werf, Lybaert ir Vanhoof (2011) išskirdami, kad neapibrėžtas algoritmas gali būti ypač svarbus nustatant tikrą procesų veikimo būdą realioje aplinkoje bei nustatant įtartinus juridinių asmenų mokėjimus.

Nepaisant didelio neapibrėžto algoritmo (angl. Fuzzy miner) populiarumo, kombinuojant šį algoritmą su kitais duomenų analizės metodais, kiti algoritmai gali būti naudingi atliekant kitas su sukčiavimo prevencija susijusias veiklas. Pavyzdžiui, euristinis algoritmas (angl. Heuristic miner) gali būti sėkmingai pritaikomas ne tik įtartinų mokėjimų nustatyme, bet ir kliento pažinimo procese nustatant įtartinumo požymius. Autoriai Silva, Tavares, Gritti ir Ceravolo (2023) nustatė, kad procesų gavybos algoritmo ir mašininio mokymosi paremtas metodas pateikia net 80% tikslumą, nustatant ar klientas turi reikšmingų sukčiavimo požymių. Toks sprendimas ilguoju laikotarpiu leistų sumažinti apkrovą kitoms prevencinėms sistemoms (pvz., mokėjimų stebėsenai), kadangi mažesnis kiekis įtartinų klientų turėtų realius šansus sėkmingai atlikti kliento pažinimo procedūrą. Taip pat procesų gavybos algoritmai gali būti taikomi nustatant sukčiavimo atvejus įprastinio audito ar viešųjų pirkimų metu. Nai et. al. (2025) atliktas tyrimas parodė, kad taikomas procesų gavybos metodas (panašus į neapibrėžtą algoritmą) gali padėti reikšmingai sumažinti korupciją viešuosiuose pirkimuose, taip

keliant vyriausybės pasitikėjimą pilietinėje visuomenėje. Tuo tarpu atliekant procesų gavybos algoritmų tinkamumo matavimus finansinio audito procesams (sukčiavimo ir klastojimo atvejams), Stephan, Lahann ir Fettke (2021) teigia, kad iš dviejų algoritmų (euristinio ir indukcinio), euristinis pateikia geresnius eksperimento metu gautus rezultatus, t.y. pritaikomumas audito procesams – 91.43%, o tikslumas aptinkant anomalijas ir įtartinus finansinius sandorius – 81.18%).

Apibendrinant galime teigti, kad procesų gavybos algoritmai, ypač neapibrėžtas (angl. Fuzzy miner), yra itin veiksmingi analizuojant didelius ir kompleksiškus mokėjimų srautus, tačiau geresnės kokybės efektyvumui pasiekti, šie metodai turi būti derinami su kitais duomenų analizės metodais. Tuo tarpu, nors indukcinio algoritmo taikymo sprendimų mokslinėje literatūroje galima rasti, tačiau didžioji dauguma tyrimų yra orientuoti į neapibrėžtą ir euristinį algoritmus, dėl jų pritaikomumo finansinių įvykių žurnalams.

Taip pat verta paminėti, kad kiekvienas iš eksperimentuose minimų algoritmų nėra visiškai tinkami įgyvendinti sukčiavimo prevenciją mokėjimų kontekste. Pavyzdžiui, nors Fuzzy miner algoritmas yra laikomas kaip vienas geriausių, galinčių puikiai susitvarkyti su triukšmu ir pateikti suprantamą modelį, tačiau Baader & Krcmar (2018) teigia, kad dėl negebėjimo atskirti pasirinkimo ir lygiagretumo, todėl gali rodyti ne visai tikslas modelio šakų kombinacijas, įterpti procesus, kurių modelyje nėra ir pan. Tuo tarpu Heuristic modelis eksperimentuose parodė savo tinkamumą triukšmingiems įvykių žurnalams ir gebėjimu filtruoti elgsenas naudojantis priklausomybių vertėmis, tačiau algoritmo veikimo kokybė priklauso ir pasirinktų specifikacijų, ir nuo pradinių duomenų parengimo, todėl galutinis rezultatas stipriai priklauso nuo žmogaus (Silva et. al., 2023). Galiausiai, atlikta literatūros apžvalga parodė, kad Inductive miner yra mažiausiai populiarus analizuojant mokėjimų įvykių žurnalus (sukčiavimo prevencijos kontekste), o to priežastis gali būti algoritmo mažesnis atsparumas triukšmui, todėl rezultate sukuriami itin kompleksiški ir sunkiai interpretuojami modeliai. Dėl panašios priežasties matome, kad eksperimentuose neregistruoja Alpha algoritmas, kadangi realioje aplinkoje algoritmas negali susitvarkyti su įvykių žurnalais, kuriuose yra neatitikimų ar klaidų.

Taigi, apibendrinus galime teigti, kad procesų gavybos metodai ir algoritmai gali būti veiksmingai pritaikomi mokėjimų stebėsenos ir finansinių nusikaltimų prevencijos kontekste, nes leidžia technologiškai iš įvykių žurnalų atkurti įvykusių procesų modelį, nustatyti nukrypimus ir identifikuoti galimus sukčiavimo atvejus. Praktikoje dažniausiai ir efektyviausiai yra taikomas neapibrėžtas algoritmas (angl. Fuzzy miner) dėl savo gebėjimo filtruoti triukšmą ir pateikti aiškiai suprantamą modelį, vengiant „spaghetti-like“ problematikos. Euristinis (angl. Heuristic miner) pasižymi savo tikslumu ir tinkamumu dirbti su triukšmingais įvykių žurnalais, todėl yra tinkamas kliento pažinimo ir auditų procesuose. Tuo tarpu indukcinis algoritmas (angl. Inductive miner) taikomas mažiausiai dėl savo jautrumo triukšmui ir sudėtingų modelių generavimo, todėl nėra

tinkamas efektyviai analizei. Atliekami eksperimentai siūlo taikyti principą, kuomet sukčiavimo prevencijos priemonė/metodas yra taikomas naudojant procesų gavybos elementą ir kitą duomenų analizės metodą/techniką, kadangi procesų gavybos algoritmai turi tam tikrą technologinį ribotumą ir reikalauja didesnio žmonių įsikišimo.

1.4 Didelių kalbos modelių (LLMs) vaidmuo ir pritaikomumas mokėjimų stebėsenoje

Kaip minėta anksčiau, mokėjimų stebėseną susiduria su itin opia problema, kuomet taisyklėmis grįstos sistemos sugeneruoja didelius kiekius klaidingai teigiamų įspėjimų, o tik maža dalis įspėjimų parodo realią sukčiavimo ar pinigų plovimo veiklą. Taip pat svarbu paminėti, kad finansų sektorius yra nuolat augantis, o didėjant elektroninių mokėjimų kiekiui, įprastinės, taisyklėmis grįstos sistemos nebegali suteikti kokybiško stebėsenos įrankio, todėl svarbu, kad sistemos būtų paremtos naujausiomis ir inovatyviausiomis technologijomis. Viena iš technologijų, galinčių pasiūlyti inovatyvesnį sprendimą yra dirbtinio intelekto technologija, kuri siūlo ir didelius kalbos modelius. Skirtingai nei įprastinės taisyklėmis grįstos sistemos, generatyvinis dirbtinis intelektas gali itin dideliu tikslumu apdoroti reglamentus, įstatymus, direktyvas, aiškinamuosius raštus ir gaires, išskiriant pagrindinius reikalavimus, suprasti ir pasiūlyti sprendimus ar sprendimų tobulinimus tiek mokėjimų stebėsenos srityje, tiek kitoms prevencijos sritims (Gurram, 2025). Dėl šių priežasčių būtina detaliau suprasti dirbtinio intelekto suteikiamą technologinį pranašumą ir detaliai analizuoti esamus ir potencialius technologijos pritaikymo metodus.

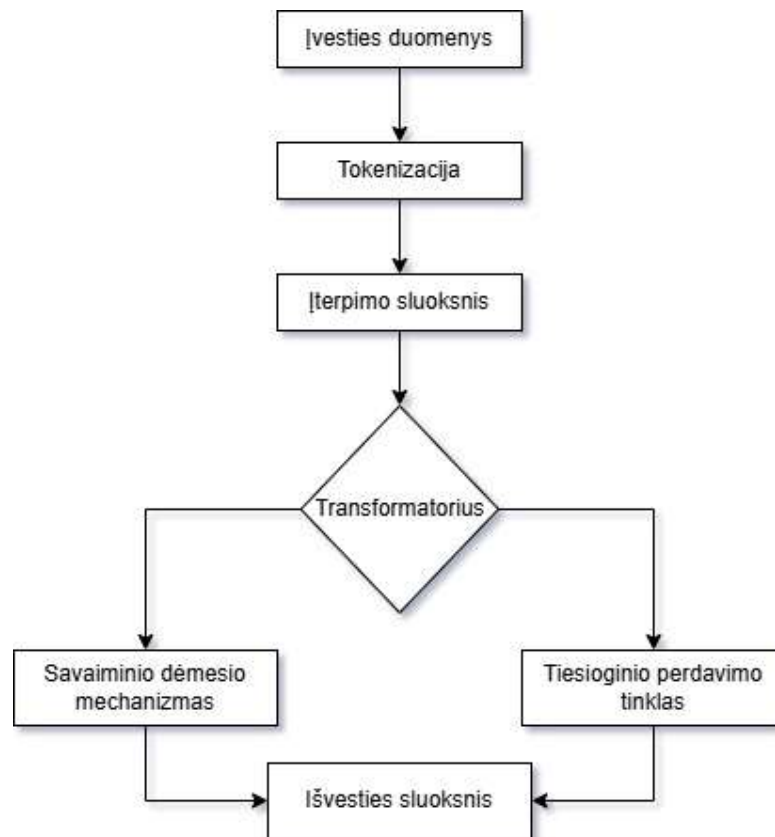
Siekiant efektyviai pritaikyti didelių kalbos modelių (angl. Large Language Models) technologiją ir funkcionalumą mokėjimų stebėsenos kontekste, būtina suprasti šios technologijos veikimo principus, veikimo žingsnius ir svarbiausius komponentus. Pagrindiniai ir svarbiausi šios technologijos veikimo žingsniai yra:

- **Tokenizacija** (angl. Tokenization) – įvesties sluoksnis skirtas apdoroti pirminę pateikiamą informaciją. Tokenizacijos veikimo principas yra paremtas teksto išskaidymu į nedalomus vienetus, vadinamus žetonais (angl. Tokens), kurie gali būti simboliai, žodžių dalys, ženklai arba žodžiai, priklausomai nuo modelio dydžio ir tipo (Naveed et. al., 2023). Tokiu principu dideli kalbos modeliai teksto pagrindu yra mokomi, kad galėtų tinkamai prognozuoti tekstą, suprasti jo konceptą ir pan.
- **Įterpimo sluoksnis** (angl. Embedding Layer) – įvestiems žodžiams ar kitai informacijai suteikiamas vektorius dimensinėje erdvėje, kuri atspindi semantinę reikšmę. Įterpimo sluoksnyje veikia teksto arba pozicinis įterpimas. Teksto įterpimai yra vektoriaus atvaizdai, kurie koduoja semantinę informaciją – toks principas plačiai naudojamas natūralios kalbos apdorojimo užduotyse (informacijos paieška, sentimentų nustatymas/atpažinimas, prekių rekomendavimas ir t.t.) (Zhang et. al., 2025). Tuo tarpu poziciniai įterpimai reikalingi, nes

transformatoriai nesupranta pateikiamų žodžių eiliškumo, todėl kontekstui priskiriamos reikšmės, kad modelis gautų teisingą kiekvieno žodžio poziciją tekste.

- **Transformatorius** (angl. Transformer Blocks) – susideda iš savaiminio dėmesio mechanizmo ir tiesioginio perdavimo tinklo. Pasiūlyta transformatoriaus infrastruktūra turi inovatyvų savaiminio dėmesio (angl. Self-attention) mechanizmą, kuriuo remiantis gali apdoroti visą įvesties seką vienu metu. Šis mechanizmas padeda kiekvienam įvesties žodžiui/žetonui priskirti reikšmingumo svorį kitam sekos elementui, todėl modelis gali lengviau įvertinti kiekvieno įvesties žodžio/žetono aktualumą kiekvieno kito žodžio atžvilgiu, efektyviau fiksuojant ilgalaikes priklausomybes ir kontekstinius ryšius (Ferraris, Audrito ir Caro, 2025). Pagal modelių veikimo principus transformatoriai gali veikti trimis skirtingais būdais – tik šifravimo, šifravimo-iššifravimo ir tik iššifravimo. Šifravimo metu apdorojama visa įvesties seka ir sukuriam kiekvienos pozicijos kodiniai atvaizdai, tuo tarpu iššifravimo metu generuojama išvesties seka. Kiekviename etape modelis apsvarto savo ankstesnius rezultatus per savaiminio dėmesio mechanizmą bei atsižvelgia į šifravimo rezultatus, taip įtraukiant reikšmingą informaciją iš įvesties sekos generuojant kitą žetoną (Ferraris et. al. 2025). Tuo tarpu tik šifravimo arba tik iššifravimo transformatoriai yra naudojami modeliuose, kurie yra skirti atlikti specifines užduotis, pavyzdžiui, tik iššifravimo transformatorius yra naudojamas sekos generavimui remiantis prieš tai pateiktu kontekstu (teksto generavimas, kalbos modeliavimas), o tik šifravimo transformatorius naudojamas kalbos supratimo užduotims atlikti (gilusis mokymasis).
- **Išvesties sluoksnis** – šio etapo rezultatas priklauso nuo naudojamo modelio tipo, pavyzdžiui, generatyviniai modeliai kaip ChatGPT yra apmokytas numatyti sekantį sekos žetoną, atsižvelgiant į ankstesnius žetonus. Tuo tarpu maskuotuose modeliuose kaip BERT ir t.t., modelis numato trūkstamus sekos žetonus. Galiausiai, nepaisant modelio tipo, galutinis išvesties taškas yra matematinė funkcija (softmax), kuri realiųjų skaičių vektorių paverčia tikimybinu pasiskirstymu.

Ši inovatyvi technologija leidžia atlikti detalesnę ir gilesnę analizę, įtraukiant ne tik sveikųjų skaičių statistinę ar priklausomybių analizę, bet įtraukiant ir normatyvinę kalbą, kintančio teksto laukus, konteksto aspektą. Tokius požymius įgalinanti technologija gali padėti sukurti ar pritaikyti naujus prevencinius modelius ar patobulinti mokėjimų stebėsenos principus, kurie atlieptų efektyvumo reikalavimą kaip dedamąją, siekiant užkirsti kelią sukčiavimui/pinigų plovimui, paraleliai užtikrinant sistemų efektyvumą.



Šaltinis: sudaryta autoriaus

pav. 4 Didelio kalbos modelio (LLM) veikimo etapų schema

Finansų ekosistemos nuolatinė raida lemia vis didėjančią ekosistemos (ir tuo tarpu mokėjimų infrastruktūrų) kompleksumą, todėl inovatyvios technologijos, tokios kaip didelių kalbų modeliai, yra testuojami, siekiant sukurti optimalesnius sprendimus, skirtus mažinti pinigų plovimą ir sukčiavimą. Bendrinė didelių kalbų modelių pritaikomumo tiek mokėjimų stebėsenoje, tiek visame finansų sektoriuje teigia, kad dirbtinio intelekto technologija turi visišką tikslumo pranašumą prieš taisyklėmis grįstas sistemas, pavyzdžiui, ypatingos svarbos raktinių žodžių kontekstinę analizę, analizuojant naujus įstatymus, direktyvas, įvairius pokyčius, kurie lemia prevencijos skyriaus darbo specifiką. Tikslumas pabrėžiamas ir įspėjimų generavimo aspekto, kuomet demonstruojami DI paremti testų rezultatai rodo klaidingų įspėjimų kiekio sumažėjimą nuo 50% iki 70% (Gurram, 2025). Kiti, žemiau tekste analizuojami moksliniai straipsniai testavimo procese renkasi įvairaus tipo didelių kalbų modelius – tiek tik iššifravimo (angl. „decoder-only“) modeliai skirti teksto generavimui ir paaiškinimui, tiek tik šifravimo (angl. „encoder-only“) modeliai, skirti teksto supratimui/kontekstinei analizei. Tokių modelių tipų naudojimas aiškinamas jų gebėjimu apdoroti tiek struktūrizuotus duomenis, tiek nestruktūrizuotus – tokia duomenų tipologija atitinka dažniausiai sutinkamus finansinių mokėjimų duomenų rinkinių požymius.

Taigi, didėjant kompleksiskumui, finansų ekosistemai reikalingi inovatyvūs sprendimai. Didelių kalbų modelių gebėjimai efektyviai analizuoti tiek struktūrizuotus, tiek nestruktūrizuotus duomenis bei gebėjimas įvertinti jų kontekstą, semantinius ryšius, yra vertinami kaip naudingi aspektai, siekiant efektyvinti tiek mokėjimų stebėsenos sistemas, tiek kitas prevencines sistemas.

Nepaisant fakto, kad didelių kalbos modelių populiarumas ir technologijos pažanga pasiekė globalaus lygio mastą tik prieš keletą metų, tačiau akademiniame lauke jau analizuojami įvairūs technologijos pritaikymo būdai mokėjimo stebėsenoje ar su pinigų plovimo prevencija susijusiuose procesuose. Pavyzdžiui, Korkanti (2024) atliekamo tyrimo metu vertina didelių kalbos modelių derinimą su įprastiniais duomenų analitikos metodais, siekiant pagerinti sukčiavimo prevencijos mechanizmų efektyvumą. Eksperimentui atlikti autorius naudoja GPT-3 modelį, o šio modelio pasirinkimas grindžiamas jo gebėjimu analizuoti tekstą, atskirti pasikartojančius veikimo modelius (angl. Patterns). GPT modelio pasirinkimas gali būti grindžiamas dėl jo technologinės specifikacijos – iššifravimo principu paremtas modelis puikiai supranta kontekstą, gali lengvai analizuoti ne tik skaitines bet ir tekstines reikšmes, todėl yra naudingas prognozuojant mokėjimo sentimentą. Tuo tarpu, naudojant anomalijų aptikimo algoritmą Isolation Forest ir didelį kalbos modelį, skirtą aptikti neįprastinius sandorius, sukurama dviejų srautų anomalijų aptikimo sistema, pagal kurios rezultatus, naudojant logistinę regresiją ir Gradient Boosting Machine formuojamas prognozavimo modelis, skirtas aptikti mokėjimus, kurie atitinka sukčiavimo kriterijus (Korkanti, 2024). Tuo tarpu Pirmorad (2025) vietoj didelio kalbos modelio naudojimo kombinuojant su tradicinėmis sistemomis, siūlo atskirą metodą, naudojant modelį kaip analitinį įrankį ir atliekant užklausų (angl. prompt) inžineriją gauti paaiškinimus/motyvus, kodėl mokėjimas turėtų būti laikomas sukčiavimo atveju. Siūloma mokėjimų duomenų rinkinius klasifikuoti į pografus, kurie susideda iš mazgų (kliento ar banko ID) ir kraštų (mokėjimo srautai). Vėliau šiuos pografus verčiant į tekstines užklausas kelių bandymų principu (angl. Few-shot) pateikiama užduotis modeliui, remiantis pagrindiniais pinigų plovimo modeliais, įvertinti ir paaiškinti ar mokėjimas turėtų būti laikomas įtartinu (įvertinus transakcijos elgesį, srauto tipą tarp mokėtojo ir gavėjo) (Pirmorad, 2025).

Kitos siūlomos sistemos ar metodai yra labiau fokusuoti į konkretų ir aiškiai užduočiai pritaikomą sprendimą, dažniausiai išnaudojant „agentinę“ didelio kalbos modelio aplinką. Panašų sprendimą siūlo Naik, Dintakurthi, Hu, Wang ir Qiu (2025) – „Co-Investigator“ – agentinė optimizuota sistema, kuri yra skirta generuoti įtartinos veiklos raportus (angl. Suspicious Activity Report). Specializuotas agentas yra skirtas apdoroti mokėjimų duomenis ir kitą KYC informaciją bei pagal atliktą apmokymą sugebėti identifikuoti plovimo tipologijas, įvertinti reguliacinius reikalavimus, suprasti naratyvus ir generuoti tekstinį rezultatą pagal sentimentą ir kontekstą. Sprendimo autoriai technologiniams reikalavimams įgyvendinti pasirenka RoBERTa modelio (tikšifravimo tipas) ir tikimybinio sekų žymėjimo modelį (CRF), kurių kombinacija suteikia aukštą

tikslumą, nuosekliai atpažinti teksto elementus, todėl toks veikimo principas padeda užtikrinti stabilų ir nuspėjamą bei optimizuotą agentinės sistemos veikimą (Naik et. al., 2025). Panašiu principu, dideliu kalbos modeliu paremtos agentinės sistemos integracijos ir veiksmingumo galimybes nagrinėja ir kiti darbai. Pavyzdžiui, Mittapelly, Tarra ir Agrahar (2024) siūloma Salesforce klientų santykių valdymo sistemą papildyti dideliu kalbos modeliu, pavyzdžiui ChatGPT, naudojanti API (liet. Aplikacijų programavimo sąsaja) prieigą bei atskiru atitikties moduliui. Siūlomas sprendimas naudoja natūralios kalbos apdorojimą (NLP) ir optinį simbolių atpažinimą (OCR), kurie apdoroja dokumentus ir kitą kliento pažinimo metu renkama informaciją, vertina riziką pagal kredito balą ir kitus kriterijus bei esant atitikties neatitikimams perduoda bylą žmogaus peržiūrai (Mittapelly et. al., 2024). Sistema yra paremtas dideliu kalbos modeliu per API sąsają ir veikia autonomiškai, tačiau galutinis sprendimas priimamas žmogaus, todėl užtikrina tinkamą kontrolę ir atitikties skaidrumą. Nors siūlomas sprendimas tiesiogiai nėra susijęs su mokėjimų stebėseną, tačiau toks techninis principas gali būti pritaikytas įtartinoms piniginėms transakcijoms stabdyti.

Taigi, didelių kalbos modelių taikymas mokėjimų stebėsenoje ar pinigų plovimo prevencijoje yra nagrinėjamas bent trimis skirtingomis kryptimis. Dažniausiai taikomas hibridinis metodas, kuomet įprastinės sistemos ar metodai yra kombinuojami su didelių kalbų modeliais, siekiant išgauti didesnę efektyvumą (mažinant klaidingų įspėjimų kiekį). Taip pat pažangūs generatyviniai modeliai tokie kaip RoBERTa gali būti pritaikomi, naudojant kaip papildoma/atskira analitinė sistema, skirta suteikti tam tikras rekomendacijas ar paaiškinimus analitikui. Galiausiai didelių kalbų modelių agentai arba šiais modeliais paremtos autonominės agentinės sistemos gali būti panaudojamos kaip efektyvūs sprendimai padedant efektyviau nustatyti įtartinas operacijas.

Kiti analitinio pobūdžio moksliniai darbai, tiriantys didelių kalbos modelių taikymą pinigų plovimo prevencijos ekosistemoje, pateikia gana didelį kiekį pavyzdžių, kaip modeliai gali būti pritaikomi darbinėje aplinkoje. Pavyzdžiui, Fan (2024) siūlo didelių kalbų modeliais paremtas aplikacijas taikyti mokėjimų segmentuose – realiojo laiko mokėjimų optimizacija, tarptautinių (SWIFT) mokėjimų ir atitikties valdyme, išmaniųjų mokėjimų asistentas (agentinė sistema) bei mokėjimų saugumo ir sukčiavimo prevencija. Nors autorius nepateikia konkrečių metodų ar modelio taikymo techninių specifikacijų, tačiau pabrėžiama, kad esami modeliai, pavyzdžiui, GPT-4 („decoder-only“), BERT („encoder-only“) bei specifinis BloombergGPT yra technologiškai pajėgūs atliepti šių dienų kompleksinių mokėjimų sistemų reikalavimus, siekiant pagerinti efektyvumą išliekant atitikties ribose. Minėtas BloombergGPT yra specifiskai finansiniams mokėjimams analizuoti skirtas modelis, kuris buvo apmokytas pasitelkiant 345 milijardus žetonų iš bendrosios paskirties finansinių duomenų rinkinių ir 363 milijardus žetonų iš Bloomberg finansinių duomenų archyvų, todėl skirtingi tyrimai fiksavo šio modelio žymų pranašumą finansinių duomenų analizės kontekste (Yan, Hu ir Zhu, 2024). Anot Yan et. al. (2024) toks modelio paruošimas konkrečiai sričiai

turi būti suprastas kaip metodologinis procesas, skirtas specifinių užduočių atlikimui (šiuo atveju finansinių duomenų analitikai), o ne BloombergGPT, autorius įvardija BERT, GPT-2 ir RoBERTa kaip didelį potencialą turinčius modelius, kurie gali revoliucionizuoti pinigų plovimo prevencijos procedūras dėl gebėjimo suprasti/interpretuoti sudėtingą finansinę kalbą ir reaguoti į atitiktis ir reguliavimo poreikius. Kiti autoriai, pavyzdžiui Han, Huang, Liu ir Towey (2020) analizuodami didelių kalbos modelių pritaikomumą pinigų plovimo prevencijos srityse išskiria ilgalaikės trumpalaikės atminties, tikimybinio sekų žymėjimo, pasikartojančių neuroninių tinklų, vardinių objektų atpažinimo ir kitus modelius/metodus, kurių taikymas gali padėti efektyviau surinkti ir atskirti svarbius finansinius duomenis, analizuoti mokėjimų srautus, suprasti jų judėjimo tendencijas, išskirti sentimentą ir pan.

Didžioji dalis analizuojamų mokslinių darbų, tiriančių didelių kalbos modelių pritaikymą mokėjimų stebėsenoje arba kitose pinigų plovimo prevencijos srityse, vienareikšmiškai sutaria šios inovatyvios technologijos teikiamas naudas. Dažniausiai išskiriama modelių gebėjimas apdoroti didelės apimties sudėtingus finansinius duomenis, itin tiksliai fiksuoti mokėjimų judėjimą, suprasti elgseną ir sentimentą bei interpretuoti duomenis, teikti išvadas ir pan. Dalis analizuojamų mokslinių darbų neapsistoja tik ties teorine analize, tačiau pateikia ir atliktus eksperimentus su aiškiais empiriniais duomenimis, todėl žemiau pateikiama lentelė su svarbiausiais atliktų eksperimentų duomenimis:

Empiriniai duomenys gauti naudojant didelių kalbos modelių technologiją finansinių nusikaltimų ir prevencijos tyrimų srityje

Autorius	Straipsnio pav.	Modelis/metodas	Finansinių duomenų rinkinys	Eksperimento rezultatai
Pirmorad, E. (2025)	<i>Exploring the In-Context Learning Capabilities of LLMs for Money Laundering Detection in Financial Graphs.</i>	Naudotas GPT-4o modelis per OpenAI API.	Naudota 2 000 mokėjimų imtis (50 % pinigų plovimo atvejai, 50 % įprastiniai mokėjimai). IBM AML sintetinis duomenų rinkinys.	63,7 proc. visų modelio gautų atvejų pateko į pasikliautino 95 proc. intervalą (su 3,3 proc. paklaida). F1 rodiklis – 65,1 proc.
Yan, Y., et. al. (2024)	<i>Leveraging Large Language Models for Enhancing Financial Compliance: A Focus on Anti-Money Laundering Applications.</i>	Naudoti BERT, GPT-2 ir RoBERTa modeliai	Simuliuoti finansinių mokėjimų duomenys iš Indijos valstybinio banko. Nurodyta imtis – 10 mokėjimų.	Visi modeliai teisingai nurodė sukčiavimo atvejus pateiktuose mokėjimų scenarijuose (F1 rodiklis – 1), nurodomas siūlymas atlikti didesnės imties eksperimentus.
Naik, P. V., et. al. (2025)	<i>Co-Investigator AI: The Rise of Agentic AI for Smarter, Trustworthy AML Compliance Narratives</i>	Naudota didelių kalbos modelių agentais grįsta architektūra, įtraukiant RoBERTa (saugumui) ir Google Gemini 2.5 Pro (nepriklausomam vertinimui).	Ekspertų įvertinti pavyzdiniai įtartinų veiklų aprašai ir kiti dokumentai skirti aprašams generuoti (vidiniai duomenys, neįvardyta).	Naratyvo generavimas pasiekė 70 proc. efektyvumą. Gautas 61 proc. darbo laiko sąnaudų rodiklis. Architektūra išskirtinai gerai išskyrė geografinės lokacijos anomalijas (100 proc.) ir sąskaitos vientisumo stebėsenos (93 proc.) rastus neatitikimus.

1 lentelės tęsinys

Korkanti, S. (2024)	<i>Enhancing Financial Fraud Detection Using LLMs and Advance Data Analytics</i>	Naudotas GPT-3 modelis, kartu su analitinėmis technikomis (anomalijų aptikimui - Isolation Forest ir Auto encoders, prognozavimui – Gradient Boosting Machine (GBM))	UCI kredito kortelių sukčiavimo atvejų duomenų rinkinys: 284 827 mokėjimai (492 sukčiavimo atvejai), taikoma SMOTE technika nesubalansuotiems duomenims.	Atliktas eksperimentas parodė dvejetainio klasifikavimo modelio veiksmingumo įverčio rezultatą tarp 0.85 ir 0.89.
Ju, C., et. al. (2025)	<i>Real-time Cross-border Payment Fraud Detection Using Temporal Graph Neural Networks: A Deep Learning Approach</i>	Naudojamas Laikinis grafų neuroninis tinklas (angl. Temporal graph neural network, TGNN).	Du duomenų rinkiniai iš finansų institucijų: A rinkinys – 2,8 milijonų mokėjimų (156 šalys) B rinkinys – 1,5 milijono mokėjimų (92 šalys)	Modelio bendras pasiektas tikslumas (per 2 skirtingus eksperimentus) yra 98,24 proc. Bendras prognozavimo efektyvumo rodiklis (F1 score) pasiekė net 0,9299. Kiti rodikliai (tikslumas, dvejetainių klasifikavimo vertinimas ir pan.) irgi peržengė 0,9 slenkstį.
T. Tang, et. al. (2024)	<i>Application of Deep Generative Models for Anomaly Detection in Complex Financial Transactions</i>	Naudojamas Generatyvinių prieštarinių tinklų (angl. Generative Adversarial Networks) ir variacinių autoenkoderių (angl. Variational Autoencoders) kombinacijos modelis.	Naudojamas sintetinis mokėjimų duomenų rinkinys PaySim.	Prognozavimo efektyvumo rodiklis (F1 score) pasiekė 0,795. Kiti rodikliai (tikslumas, informacijos išgavimas ir t.t.) varijavo tarp 0,763 ir 0,946 bei aplenkė kitus testuotus modelius (GAN, VAE, GAT ir pan.).

Šaltinis: sudaryta autoriaus

1 lentelėje pateikiami susisteminti atliktų eksperimentų empiriniai duomenys leidžia teigti, kad didelių kalbos modelių taikymas tiek mokėjimų stebėsenoje, tiek kitose pinigų plovimo prevencijos srityse leidžia pasiekti reikšmingų efektyvumo rezultatų. Gauti rezultatai rodo, kad didelių kalbos modelių taikymas gali leisti pasiekti prognozavimo efektyvumą režiuose nuo 0,65 iki 1, sumažinti laiko sąnaudas, pavyzdžiui, iki 61 proc., bei išsiskiria gebėjimu aptikti anomalijas, elgsenos vientisumą, klasifikuoti ir interpretuoti duomenis. Dėl riboto kiekio atliktų eksperimentų taikant didelius kalbos modelius siekiant efektyviau nustatyti įtartinus mokėjimus ar gerinti pinigų plovimo prevencijos procedūras, 1 lentelėje pateikiami dviejų eksperimentų rezultatai, kuriuose buvo naudojami kiti dirbtinio intelekto modeliai/metodai. Šių eksperimentų rezultatai indikuoja prognozavimo efektyvumo režius tarp 0,795 ir 0,9299, o tuo tarpu kiti rodikliai (tikslumas ir efektyvumas), variavo tarp 0,763 ir 0,946, kas leidžia teigti, kad dirbtinio intelekto technologija gali reikšmingai daryti įtaką mokėjimų stebėsenos ir kitų prevencijos procedūrų efektyvumui. Tam antrina ir Olateju et. al. (2024), kuris analizuodamas 50 skirtingų mokslinių darbų nustatė, kad vidutiniškai dirbtiniu intelektu paremti sprendimai (nebūtinai dideli kalbos modeliai) gali sumažinti klaidingai teigiamų įspėjimų kiekį vidutiniškai iki 4,23 proc. (0,41 proc. paklaida) bei indikuoja aukštą tikslumą siekiant nustatyti anomalijas, įprastiems mokėjimams nebūdingus bruožus ir t.t. Apibendrinus, galime teigti, kad tiek dirbtinio intelekto technologija (ir skirtingi modeliai), tiek išskirtinai didelių kalbos modelių pritaikomumas įvairiuose pinigų plovimo prevencijos procesuose turi reikšmingą empirinę įtaką, ypač laiko resursų taupyme, anomalijų ir konteksto aptikime bei semantinės informacijos analizavime.

Moksliniai darbai tiriantys didelių kalbos modelių taikymą mokėjimų sistemose ar pinigų plovimo prevencijos etapuose įvardija šios inovacijos teikiamus pranašumus bei kylančius trūkumus. Pavyzdžiui, Naik et. al. (2025) pabrėžia, kad agentinės didelių kalbos modelių taikymas pinigų plovimo prevencijos procesuose žymiai didina analizės greitį ir darbo eigos efektyvumą. Efektyvumo svarbą pabrėžia ir anksčiau pateikti susisteminti tyrimų empiriniai rezultatai, rodantys gerėjančius efektyvumo ir taiklumo rodiklius. Taip pat didelė dalis autorių sutaria, kad vieni iš ryškiausių didelių kalbos modelių aspektų yra jų gebėjimas suprasti duomenų sentimentą, gebėjimas analizuoti kompleksiškus finansinius duomenis, sugebėjimas sekti besikeičiančią mokėjimų elgseną ir teikti išsamius aprašymus/aiškinimus, suteikiant kontekstą. Tačiau, anot Pirmorad (2025), nors nagrinėjami modeliai yra perspektyvūs, tačiau praktikoje kyla įvairių apribojimų, pavyzdžiui, didelių kalbos modelių skaičiavimo sąnaudos (teikiant išvadas, paaiškinamuosius tekstus ir pan.) yra per daug didelės, kad būtų galima tokį sprendimą pritaikyti didelio intensyvumo aplinkose. Anot Fan (2025) taip pat pastebi, kad dideli kalbos modeliai kelia iššūkių dėl duomenų apsaugos, algoritmo šališkumo problematikos, modelio interpretacijų ir haliucinacijų rizikos. Modeliui veikiant itin kompleksiskame finansinių duomenų tinkle, tokios klaidos gali būti reikšmingos tiek finansų institucijos, tiek

klientams. Dėl šių išvardintų problemų kyla per daug didelė rizika, kad mokėjimų stebėsenos ir kitų prevencijos procedūros bus atliekamos netaisyklingai, neatsižvelgiant į įstatymų ir reikalavimų kontekstą, remiantis haliucinacijomis ir pan., kas gali lemti didelę finansinę žalą finansų įstaigai, taikančiai pažangius metodus (reputacinę ir reguliacinę rizikos). Todėl, nors technologinis pažangumas ir gerina resursų išteklių panaudojimą, tačiau kvalifikuotas specialistas privalo tiek prižiūrėti modelio darbą, tiek priimti galutinį sprendimą.

Taigi, apibendrinant galime teigti, kad didelių kalbos modelių technologijos taikymas mokėjimų stebėsenoje bei kituose pinigų plovimo prevencijos procesuose gali suteikti reikšmingą efektyvumo pokytį, siekiant užtikrinti didesnę analizės tikslumą ir efektyvumą bei reikšmingai sumažinti klaidingų įspėjimų kiekį. Kuomet įprastinės taisyklėmis grįstos sistemos dažniausiai dėl savo riboto prisitaikymo prie dinaminės finansų sektoriaus aplinkos dažnai generuoja perteklinį klaidingų signalų kiekį, didelių kalbų modelių technologija gali būti panaudota kaip inovatyvus proveržis, kuriant pažangesnį rizikų valdymo mechanizmą mokėjimų stebėsenoje. Remiantis atliktų tyrimų empiriniais duomenimis, didelių kalbų modeliai, tokie kaip GPT, BERT, RoBERTa ar BloombergGPT demonstruoja gana aukštą prognozavimo efektyvumą (F1 score rodiklis varijuoja nuo 0,65 iki 1,00), geba analizuoti tiek struktūrizuotus, tiek nestruktūrizuotus duomenis, įvertinti semantinį kontekstą bei generuoti paaiškinimus ir rekomendacijas, padedančias analitikams priimti pagrįstus sprendimus. Panašius rezultatus generuoja ir kiti dirbtinio intelekto modeliai (grafiniai neuro-tinklai, Gradient Boosting Machine ir kt.), kurie neretai būna derinami kartu su didelių kalbų modeliais. Vis dėlto, šių technologijų praktinis taikymas susiduria su įvairiais iššūkiais – didelės skaičiavimo sąnaudos, duomenų saugumas/privatumas, algoritmo šališkumas, galimos modelių haliucinacijos ir kt. – todėl inovatyvios technologijos taikymas privalo būti prižiūrimas kompetentingų specialistų, kurie galėtų užtikrinti rezultatų patikimumą ir atitiktį reguliaciniams reikalavimams. Nepaisant to, galima konstatuoti, kad didelių kalbų modelių technologija turi pakankamą potencialą transformuoti mokėjimų stebėsenos sistemų principą, didinant jų tikslumą, efektyvumą bei geresnį pritaikomumą nuolat kintančiai aplinkai, t.y. sudėtingėjant mokėjimų srautų struktūroms, griežtėjant atitikties reikalavimams ir pan.

2. SIŪLOMO METODO SKYRIUS

Antroje darbo dalyje pateikiamas išsamus siūlomo metodo aprašas, įvardijant esmines metodo dalis, veikimo principus, naudojamus duomenis ir kt. Siekiant spręsti teigiamai klaidingų įspėjimų kiekį mokėjimų sistemose, siūlomas procesų gavybos ir didelių kalbų modeliais grįstas hibridinis metodas, skirtas tiksliau ir efektyviau nustatyti sukčiavimo atvejus didelės apimties įvykių žurnale. Šioje dalyje taip pat pateikiama bendra metodo struktūra (su detalio veikimo schema), procesų gavybos ir didelių kalbos modelių panaudojimas, požymių sintezė ir kiti aspektai, susiję su modelio veiksmingumo pagrįstumu, kuris padės spręsti gana aktualią problemą, su kuria susiduria bankai, finansų institucijos ir kitos įstaigos, siūlančios finansinius produktus.

2.1 Siūlomo metodo konceptas ir pagrįstumas

Kaip minėta ankstesniame skyriuje, pagrindinė darbe keliama problema yra didelis kiekis klaidingai teigiamų įspėjimų, kurie sukuria papildomus kaštus finansų įstaigoms, lemia neefektyvų personalo darbo jėgos paskirstymą bei lemia neigiamą kliento patirtį naudojantis finansiniais produktais, todėl galiausiai patiriama ir reputacinė žala. Viena iš esminių šios problemos priežasčių yra dabartinių mokėjimų stebėsenos sistemų technologinis pagrįstumas automatiškai nekintančiomis taisyklėmis, kurios yra formuluojamos pagal esamus įstatymus ir kitus reguliacinius reikalavimus. Įprastinė taisyklėmis grįsta sistema atsilieka nuo finansinių produktų inovacijų plėtros greičio ir sukuria itin gerą terpę asmenims, užsiimantiems sukčiavimu ar pinigų plovimu. Taip pat tokia sistema technologiškai nėra pritaikyti analizuoti finansinių įvykių, sunku pastebėti anomalijas ar įtartinos veiklos požymius bei taisyklių bendrą efektyvumą.

Dėl įprastinių sistemų išvardintų trūkumų galime suprasti aiškią paklausą naujiems metodams ar produktams, kurie sugebėtų tinkamai pakeisti pasenusią taisyklėmis grįstą sistemos principingumą. Atlikta literatūros analizė parodo, kad moksliniu lygiu yra atlikta daug įvairių tyrimų naudojant įvairius duomenų ir procesų analizės algoritmus bei kitas inovacijas, pavyzdžiui, dirbtinį intelektą. Taip pat dažniausiai siūlomi metodai yra taikomi hibridiniu principu, pavyzdžiui, siūlomas procesų gavybos algoritmo ar didelių kalbų modelio pritaikymas yra paremtas kitais pažengusiais duomenų analizės ir apdorojimo metodais (mašininis mokymasis, neuro-tinklų taikymas, analitinės technikos ir t.t.). Atsižvelgiant į susistemintą informaciją literatūros apžvalgos metu, siūlomam metodui pasirenkamas procesų gavybos neapibrėžto algoritmas (angl. Fuzzy miner) ir GPT modelis (gpt-4.1-mini). Fuzzy miner pasirenkamas dėl jo galimybių apdoroti didelės apimties triukšmingą finansinių įvykių žurnalą bei sugeba generuoti aiškiai suprantamą veiksmų modelį. Tuo tarpu GPT didelis kalbos modelis pasirenkamas atsižvelgiant į kitus literatūrinius darbus, analizuojančius ir vertinančius jo efektyvumą ir pritaikomumą siekiant analizuoti ir tiksliai identifikuoti sukčiavimo ir pinigų

plovimo atvejus. Taip pat atsižvelgiama į modelio pajėgumus ir prieinamą API prieigą, kuri leidžia efektyviau ir griežčiau valdyti modelio parametrus, taip didinant modelio darbo efektyvumą ir gaunamų rezultatų prognozavimą.

Naudojant hibridinį metodą (kombinuojant procesų gavybos algoritmą ir didelį kalbos modelį) siekiama sistemingai analizuoti skirtingų kintamųjų sąsajas su anomalijomis (pinigų plovimo atvejais). Remiantis pasiūlytu modeliu yra generuojamos naujos savybės (angl. Features) gautos iš procesų gavybos algoritmo bei didelio kalbos modelio (taikant užklausų (angl. Prompt) inžineriją). Išgautos naujos savybės, tikimasi, kad pateiks naujus požymius, galinčius padėti efektyviau nustatyti pinigų plovimo įvykius, kurių negali nustatyti įprastinė taisyklėmis paremta mokėjimų stebėsenos sistema. Taip pat būtina paminėti, kad metodas kuriamas dedukciniu principu, apibendrinus Lietuvos ir užsienio šalių autorių darbus bei esamus metodus/produktus.

2.2 Bendroji metodo struktūra

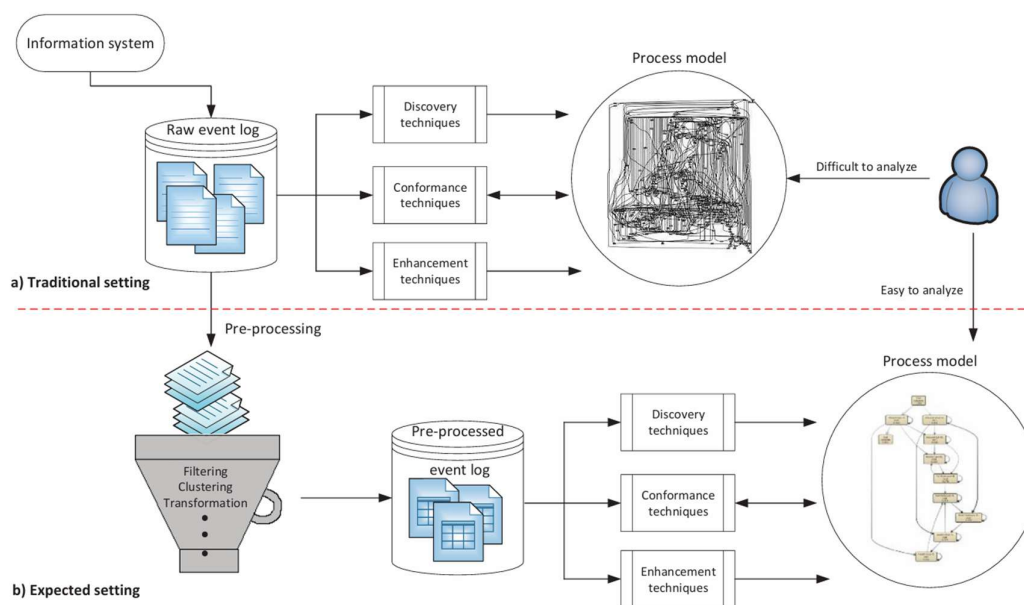
Siekiant geriau suprasti ir pateikti siūlomo metodo veikimo principą, būtina suprantamai pateikti metodo veikimo struktūrą, aprašant esminius veikimo žingsnius ir sąsajas tarp naudojamų komponentų. Todėl, šiame poskyryje pateikiamas detalus pagrindinių metodo žingsnių principingumas, naudojamų algoritmų ir modelių technologinės specifikacijos bei vizualinė veikimo schema. Taip pat pateikiamas detalus duomenų parengimas ir naudojamų kintamųjų svarba, siekiant metodo metu išgauti tinkamas naujas savybes (angl. Features) reikalingas modelio veiksmingumui pagrįsti.

2.2.1 Duomenų rinkinio paruošimas

Siūlomas metodas yra pagrįstas hibridiniu modeliu naudojant tiek procesų gavybos algoritmą, tiek didelį kalbos modelį, todėl siekiant užtikrinti efektyvų ir naudingą šių algoritmų ir modelių taikymą, būtina tinkamai paruošti turimus mokėjimų duomenų rinkinius. Šis etapas yra svarbus, siekiant, kad procesų gavybos algoritmas pateiktų kuo aiškesnį ir neiškraipytą proceso modelį, o didelių kalbu modelis pateiktų tikslesnes semantinių duomenų ir kitų rodiklių interpretacijas – taip gaunami tikslesni rezultatai.

Praktikoje duomenų išankstinis apdorojimas yra nedaug dėmesio susilaukiantis aspektas, tačiau yra labai svarbus ir galintis ženkliai lemti galutinius rezultatus ir išvadas. Pavyzdžiui, Garcia, Luengo ir Herrera (2015) atkreipia dėmesį, kad išankstinio apdorojimo etapas yra itin laisvas ir kiekvieno žmogaus skirtingai interpretuojamas procesas, todėl tam tikros žmogiškosios klaidos gali lemti nelogiškas ar tikrovės neatitinkančias duomenų sekas, duomenų rinkiniuose gali likti daug tuščių (angl. Null) reikšmių ar perteklinės informacijos, kuri daro neigiamą įtaką tolimesniam procesų gavybos algoritmui ir galutinių rezultatų bei išvadų formavime. Marin-Castro ir Tello-Leal (2021)

taip pat antrina, kad įvykių žurnalas yra svarbiausia procesų gavybos įvestis, todėl išankstinis duomenų (įvykių) apdorojimas, t.y. trūkstamų, klaidingų, triukšmingų ar pasikartojančių reikšmių šalinimas leidžia išvengti „spaghetti-like“ modelio problematikos bei parodo aiškiau suprantamą veiksmų proceso modelį tolimesniam analitiniam procesui, kaip pavaizduota 5 paveiksle:



Šaltinis: Marin-Castro, H. M. ir Tello-Leal, E. (2021). *Event log preprocessing for process mining: a review*, p. 2.

pav. 5 Išankstinio duomenų apdorojimo įtaka procesų gavybos modelio sudarymui

Atsižvelgiant į pirminių duomenų kokybės svarbą procesų gavybos atlikimui, suprantame, kad nėra vieningo būdo kaip tinkamai paruošti ir apdoroti įvykių žurnalą, kadangi kiekvienam tiriamam atvejui yra būtinas konteksto, keliamų problemų, numatomų sprendimo būdų ir pan., įvertinimas. Tačiau Konkarti (2024) teigia, kad išankstiniam duomenų parengimui galima taikyti daug skirtingų technikų ir jas kombinuoti pagal kiekvieno atvejo reikalavimus, pavyzdžiui, duomenų normalizavimas, nesubalansuotų duomenų tvarkymas ar duomenų suskirstymas į mokymo ir testavimo rinkinius (labiau taikomi didelių kalbų modeliams ar kitiems DI modeliams). Tuo tarpu procesų gavybos srityje, Garcia et. al. (2015) pažymi, kad išankstinių duomenų apdorojimas gali būti atliekamas įvairiomis formomis, tokiomis kaip duomenų transformacija, duomenų integracija, duomenų normalizacija, trūkstamų duomenų papildymas bei triukšmo identifikavimas ir šalinimas. Šie išvardinti veiksmai padeda efektyviau atlikti pirminį duomenų rinkinio paruošimą bei užtikrina, kad gaunami rezultatai bus geresnės kokybės nei naudojant neapdorotus duomenis.

Siekiant išgauti gerus semantinės analizės rezultatus iš didelio kalbos modelio (GPT API), būtina tinkamai paruošti tekstinę informaciją modelio interpretacijoms ir vertinimui gauti. Prieš

modeliui pateikiant duomenis interpretacijai ir vertinimo rezultatams gauti, galima pritaikyti aukščiau minėtas išankstinio duomenų paruošimo technikas, kurios gali padėti gauti geresnius semantinės analizės rezultatus. Pavyzdžiui, atrenkant duomenų rinkinio fragmentus (laisvo teksto formas ar kitą tekstinę informaciją), būtina įvertinti duomenų vientisumą, pasikartojančios informacijos tekste dažnumą ir t.t. Taip pat, ruošiant tekstinės informacijos rinkinį didelio kalbos modelio semantiniam vertinimui, verta atkreipti dėmesį į galimybę konvertuoti specifinių skaitinių kintamųjų reikšmes į tekstinę, suskirstant skaitinę informaciją į tekstines kategorijas (pagal atrinktus skaitinius režius). Toks metodo taikymas, leidžia suteikti papildomą kontekstą dideliame kalbos modeliui, todėl grąžinamos reikšmės ir paaiškinamoji informacija tampa labiau tikslesnė.

Atliekant išankstinį duomenų paruošimą pagal siūlomo metodo principą išskiriami du aspektai, pagal kuriuos yra ruošiamas duomenų (įvykių) rinkinys – vienas skirtas procesų gavybos algoritmo analizei, o kitas – didžiojo kalbos modelio (GPT API) analizei. Pirmiausiai suformuotame įvykių žurnale išskiriamas unikalūs proceso atvejis (caseID), veiklos tipas ir įvykio laikas, o siekiant užtikrinti globalesnę analizės prizmę, išskiriami papildomi laiko dimensijų atributai (pvz., mėnesio ar ketvirčio intervalai) bei papildomi filtravimui skirti vartotojų atributai, kombinuojant su esama sukčiavimo žyma. Tuo tarpu dideliame kalbos modeliui (GPT) iš atrinktų tekstinių ir kontekstinių atributų sukuriamas semantinis operacijos naratyvas, kuris perteikia kiekvieno įvykio esmę natūralia kalba. Papildomai, atrinkti skaitiniai atributai yra konvertuojami į tekstines kategorijas pagal nustatytus skaitinius režius, todėl suteikiamas papildomas kontekstas detalesnėms interpretacijoms. Režiumuojant, šiame etape parengiami du duomenų pogrupiai, skirti tolimesnei analizei, siekiant nustatyti bendrą požymių rinkinį, pagal kurį atliekama tolimesnė veiksmingumo analizė, nustatant įtartinas mokėjimų operacijas.

2.2.2 Procesų gavybos (Fuzzy Miner) taikymas

Atlikus išankstinį duomenų rinkinio paruošimą, sekantis svarbus metodo etapas yra procesų gavybos atlikimas. Šis procesas gali rekonstruoti realų mokėjimų vykdymo procesą ir nustatyti elgsenos modelius, kurie gali parodyti nukrypimus nuo įprastinės veiklos, indikuojant įtartiną ar sukčiavimo veiklą. Atliekant procesų gavybos analizę, gautas modelis vizualiai pateikia identifikuotų veiklų sekas, jų pasikartojimus bei netipinius finansinių mokėjimų elgsenos modelius – svarbi informacija, galinti teisingai parodyti pinigų plovimo ar kitokio pobūdžio nusikaltimų atvejus.

Vertinant finansinių operacijų duomenų rinkinius pastebima, kad tokio tipo rinkiniai pasižymi dideliu ir triukšmingu įrašų kiekiu, kadangi pati finansų rinka yra itin kompleksiška, įvykiai yra nehomogeniški ir dinamiški, t.y. veiksmų sekos dažnis nėra vientisus, o saugomi duomenys dažniausiai yra skirtingi. Siekiant atlikti tinkamą procesų vertinimą ir atsižvelgiant į turimų duomenų kompleksiskumą, siūlomo metodo procesų gavybos modulyje naudojamas neapibrėžtas procesų

gavybos algoritmas (angl. Fuzzy Miner). Siekiant tinkamai pritaikyti minimą algoritmą siūlomam metodui, būtina detaliai suprasti algoritmo funkcionalumą ir veikimo principus.

Baader ir Kremer (2018) akcentuoja, kad Fuzzy miner algoritmas iš anksto nestruktūrizuoja duomenų, bet naudoja teminės kartografijos metodus, kuomet pagrindiniu tikslu tampa detalių abstrahavimas ir dažniausiai pasikartojančių proceso atvejų vizualizacija. Algoritmas matuoja procesų svarbą pagal pasikartojimo dažnumą ir eiliškumą bei procesų koreliaciją, todėl formuojamame modelyje išskiriamus pagrindinius procesus galima lengviau analizuoti bei daryti atitinkamas išvadas. Toks funkcionalumas analitikui leidžia įvertinti bendrą ir įprastą veiksmų logiką, esančią analizuojamame įvykių rinkinyje, tačiau tuo pačiu filtravimo mechanizmai leidžia detaliau analizuoti mažiau reikšmingus įvykius, įvairius nukrypimus bei neįprastą veiklą. Taip pat verta paminėti, kad Fuzzy miner algoritmas nereikšmingas ir ribines reikšmes grupuoja į bendrą vienetą, todėl naudojantis filtravimo mechanizmu galima įprastinę veiklą palyginti su nukrypimais, ieškant savybių, leidžiančių paaiškinti nukrypimo priežastis ir bendrus nukrypimų bruožus. Bendrinis šio algoritmo funkcionalumas pritaikytas finansinių mokėjimų įvykių rinkinyje leidžia identifikuoti ne tik įprastinę mokėjimų veiklą, bet tuo pačiu ir filtruoti anomalijas ar ribines reikšmes, galinčias parodyti sukčiavimo ir nelegalios veiklos įvykius.

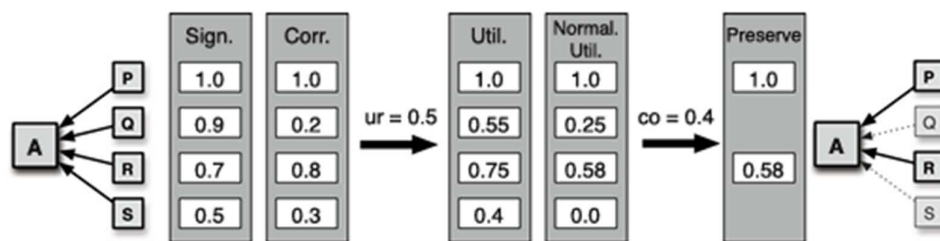
Vertinant Fuzzy Miner algoritmo funkcionalumą suprantame, kad šis algoritmas yra skirtas apdoroti didelius ir nestruktūruotus įvykių žurnalus, kuriuos apdorojantys kiti procesų gavybos algoritmai susiduria su „spaghetti-like“ problematika (modelį sunku interpretuoti dėl jo kompleksiško). Gunther ir Aalst (2007) darbe ne tik aprašo Fuzzy miner algoritmo veikimo principus ir esmę, bet ir konceptualiai pateikia kiekvieno iš svarbiausių rodiklių skaičiavimo matematinę logiką, kurios pateikiamos žemiau:

- **Reikšmingumas** (angl. Significance) – nusako veiklos ar veiklos porų svarbą pagal tai, kaip dažnai jos pasirodo įvykių žurnale. Anot autorių (Gunther ir Aalst, 2007), matematiniu principu rodiklis išvedamas santykiu vertinant kaip dažnai įvykis A pasirodo, lyginant su visu kiekiu įvykių esančių žurnale. Analogiškai teigiama, kad veiklų perėjimas (pvz., iš A į B) turi dvejetainę reikšmę (angl. Binary significance), kuri parodo šio veiksmų perėjimo santykinį dažnumą lyginant su kitų veiksmų perėjimais modelyje. Taip pat nurodomi ir kiti galimi reikšmingumo aspektai, kaip įvykio trukmės reikšmė ar maršruto reikšmė, o galutiniame rezultate turimos reikšmės yra normalizuojamos į intervalinę reikšmę nuo 0 iki 1, pagal kurią vertinami proceso įvykių tarpusavio reikšmingumai (Gunther ir van der Aalst, 2007).
- **Koreliacija** (angl. Correlation) – skaičiuojamas tarp dviejų iš eilės einančių įvykių, nustatant kaip glaudžiai įvykiai yra susiję bendriniame kontekste. Autoriai (Gunther ir van der Aalst, 2007) nurodo, kad koreliacija galima vertinti pagal kelių rūšių požymius, pavyzdžiui, laiko artumą, įvykio iniciatorių ir duomenų panašumą, vertinant kaip įvykių kontekstai glaudžiai

yra susiję. Kaip ir reikšmingumo rodiklis, koreliacijos vertė taip pat yra normalizuojama intervale tarp 0 ir 1 (kuo arčiau vieneto reikšmės, tuo koreliacija didesnė).

Kuomet remiamasis dviem rodikliais – reikšmingumu ir koreliacija – algoritmas gali dinamiškai supaprastinti procesą, išlaikydamas tik svarbiausius ir tarpusavyje glaudžiai susijusius elgesio modelius, o mažiau reikšmingas veiklas grupuoja ar paslepia (Jans et. al., 2011).

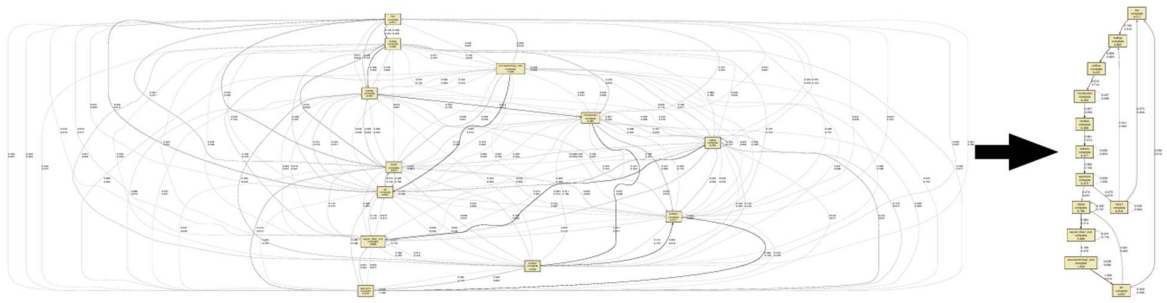
Kitas svarbus neapibrėžto algoritmo (angl. Fuzzy miner) techninis komponentas, apie kurį užsimenama aukščiau yra ribinių reikšmių filtravimas (angl. Edge filtering), kuris leidžia kontroliuoti modelio detalumo apimtį, išskirti tik esmines modelio procesų dalis ar analizuoti kraštines reikšmes. Gunther ir van der Aalst (2007) nurodo, kad po reikšmingumo ir koreliacijos rodiklių apskaičiavimo kiekvienai veiklai ir jos ryšiams, algoritmas taiko filtravimo slenksčius (angl. Thresholds), pagal kuriuos galime apibrėžti, kokie procesai bus atvaizduoti galutiniame modelyje, taip išfiltruojant ribines reikšmes. Autoriai nurodo, kad tiek ribiniai procesai, tiek svarbūs procesai yra apskaičiuojami naudojant funkciją $util(A, B) = ur \times sig(A, B) + (1 - ur) \times cor(A, B)$, kurioje kintamasis ur reprezentuoja vartotojo nustatytą svorį. Praktikoje tokios funkcijos pritaikomumą galime suprasti pagal žemiau pateiktą 6 paveikslą:



Šaltinis: Gunther, C. W., ir Aalst, van der, W. M. P. (2007). *Fuzzy mining- adaptive process simplification based on multi-perspective metrics*, p. 339

pav. 6 Įeinančių briaunų į mazgą A filtravimo procesas Fuzzy Miner algoritme

6 paveiksle patektas įvykis A turi keturias galimas jungtis su kitų procesų dalimis. Algoritmui apskaičiavus kiekvienų įvykių tarpusavio reikšmingumą ir koreliaciją, naudojant vartotojo nustatytą svorį, apskaičiuojama ir normalizuojama naudos rodiklis, kuris nustato ribines reikšmes ir išfiltruoja jas iš pagrindinio modelio. Šio Fuzzy Miner algoritmo funkcionalumo svarbą autoriai Gunther ir van der Aalst (2007) parodo, pateikdami procesų gavybos metu sukurtų modelių palyginimą 7 paveiksle, kuomet kairiajam modeliui nebuvo taikomas ribinių reikšmių filtravimas, o dešinėje pusėje - buvo

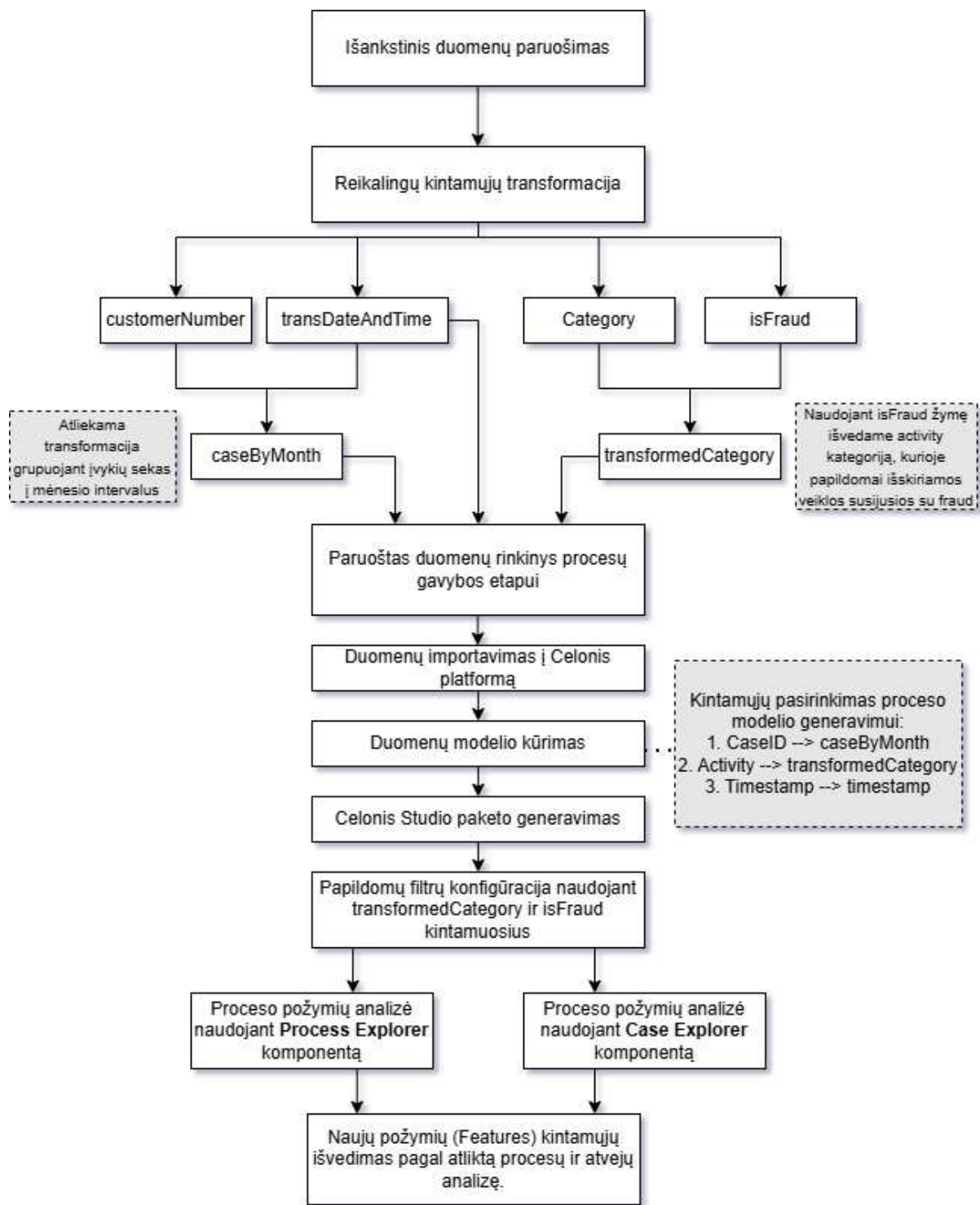


Šaltinis: Gunther, C. W., ir Aalst, van der, W. M. P. (2007). *Fuzzy mining- adaptive process simplification based on multi-perspective metrics*, p. 339

pav. 7 Procesų modelio pavyzdys netaikant ribinių reikšmių filtravimo (kairėje) ir taikant ribinių kraštinių filtravimą (dešinėje)

Šis algoritmo (Fuzzy Miner) funkcionalumas laikomas vienu iš didžiausių pranašumų, lyginant su kitais pagrindiniais algoritmais, kas leidžia daug efektyviau išgauti tiek pagrindinius procesų įvykius, tiek nustatyti ribinius atvejus bei įvertinti kaip jie skiriasi. „Spaghetti-like“ modelio problematikos išsprendimas šiuo atveju leidžia analitikui greičiau ir efektyviau suprasti esmines procesų modelio detales ir bendrą vaizdą, o filtravimo funkcijos parametrų redagavimas leidžia filtruoti ribines reikšmes ir kitas anomalijas, kurios mokėjimų kontekste gali parodyti sukčiavimo ir pinigų plovimo atvejus.

Įvertinę neapibrėžto algoritmo (Fuzzy miner) technines specifikacijas ir veikimo principą bei remiantis aukščiau atlikta literatūros analize, galime teigti, kad šis algoritmas gali būti panaudotas pritaikant procesų gavybos eigą siūlomame metode. Sudarytame šio darbo metode, atliekant procesų gavybos modelio formavimą ir analizę, siūloma naudoti viešai prieinamą komercinį Celonis procesų gavybos įrankį, kurio veikimo principas yra paremtas vidiniu Procesų atradimo varikliu (angl. Process Discovery Engine). Nepaisant to, kad Celonis vidinė dokumentacija nėra viešai prieinama, tačiau įvairūs autoriai, tokie kaip Baader ir Krcmar (2018), Celik ir Akcetin (2018) ir kiti pabrėžia, kad Celonis platformos veikimo principas yra paremtas neapibrėžtu (Fuzzy Miner) algoritmo logika. Įvertinus Celonis suteikiamą vartotojo sąsają ir suteikiamus modelio elgsenos filtrus bei parametrų reguliavimo galimybę, galime sudaryti procesų gavybos veiksmų schemą, kuri apibūdina visus siūlomo metodo procesų gavybos modulio atliekamus etapus, siekiant išvesti naujus požymius (features), reikalingus tiksliai identifikuoti įtartinus mokėjimus. Minima detali procesų gavybos veiksmų schema yra pateikiama 8 paveiksle:



Šaltinis: sudaryta autoriaus

pav. 8 Siūlomo metodo procesų gavybos modulio veikimo schema

Atliekant išankstinį duomenų apdorojimą užtikriname, kad visi reikalingi procesų elementai yra tinkamai paruošti, siekiant sudaryti kuo tikslesnę ir aiškesnę procesų modelį tolimesnei analizei.

Šis procesas yra svarbus, kadangi Celonis platformoje procesų gavyba priklauso nuo trijų svarbių elementų – bylos numeris (angl. Case ID), veikla (angl. Activity) ir laiko žyma (angl. Timestamp) – kurie lemia galutinio modelio kokybę. Kadangi finansinių mokėjimų įvykių sekos dažniausiai atvaizduoja vieną įvykį (mokėjimo išsiuntimą ar gavimą), todėl siekiant atlikti aukštos kokybės procesų gavybos eigą, iš esamų laiko žymų formuojame globalius kategorinius laiko masyvus grupuojant dieninius mokėjimus į vieno mėnesio ir ketvirtinius (3 mėnesių) rinkinius. Atliekant šį veiksmą atliekame kliento identifikacinio numerio (*customerNumber*) ir mokėjimo įvykdymo laiko žymos (*transDateAndMonth*) kintamųjų transformaciją, išvesdami naują kintamą – įvykis pagal mėnesį (*caseByMonth*). Toks duomenų koregavimas yra svarbus šiuo atveju, kadangi finansinio sukčiavimo ir pinigų plovimo procesai yra imlūs laiko masyvui dėl įvairių perlaidų strategijų ir struktūrų (pvz., nešvarių pinigų įliejimo ir sluoksniavimo etapai). Taip pat sukčiavimo ir pinigų plovimo prevencijos analizę sunku atlikti kokybiškai analizuojant vieną įvykį, todėl taikomas laiko retrospektyvos grupavimas leidžia efektyviau įvertinti finansinių įvykių ypatybes ir elgseną. Galiausiai, svarbu paminėti, kad atliekant modelio analizę turime galimybę filtruoti procesus tiek pagal tarpusavio veiksmų reikšmingumą/koreliaciją, tiek analizei naudojant kitus kintamuosius ir ribinių reikšmių filtrus (specifinių modelio etapų analizė). Dėl šios priežasties išankstinio duomenų apdorojimo metu sukuriame papildomą kintamą – *transformedCategory* – kuris susideda iš *Category* ir *isFraud* kintamųjų. Tokia kombinacija leidžia atskirti ir filtruoti procesų modelį į dvi dalis – įprastinė finansinė veikla ir finansinė veikla su sukčiavimo atvejais, todėl galima efektyviau analizuoti sukčiavimo įvykių veiklą, jos atsiradimą įprastiniuose finansiniuose mokėjimuose. Šis filtravimo ypatumas atskiriant veiklas į atskirus segmentus leidžia analitikui išskirti naujus požymius, susijusius su sukčiavimo ir pinigų plovimo elgsenos atvejais.

Atlikus sudaryto procesų modelio analizę ir pasitelkiant filtravimo funkciją pagal kitus esamus ar sukurtus kintamuosius, šiame metodo etape išvedame naujus procesinius požymius (angl. Features), kuriuos laikome kaip reikšmingus, siekiant efektyviau nustatyti finansinio sukčiavimo ir pinigų plovimo atvejus stebėsenos metu. Šiame metodo etape siekiama išgauti naujus požymius, kurie gali būti naudojami analizėje nustatant sukčiavimo ir pinigų plovimo atvejus finansinių mokėjimų kontekste. Sudarytame procesų modelyje galime identifikuoti tiek įprastinius mokėjimų procesus, tiek nukrypimus, kurie praktikoje gali padėti nustatyti finansinio sukčiavimo ar pinigų plovimo atvejus, kurie įsimašo į visą didelį finansinių įvykių srautą.

2.2.3 Didelio kalbos modelio taikymas

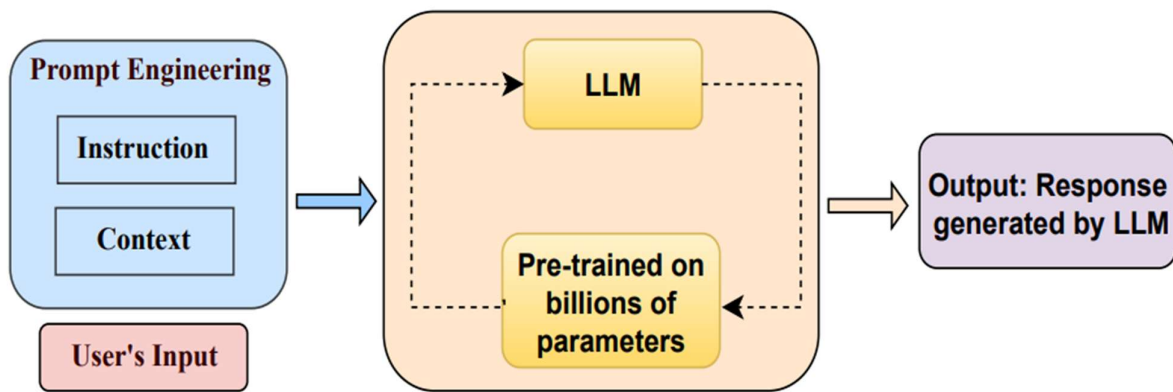
Šiame metodo etape yra siūloma naudoti didelio kalbos modelio funkcionalumą, kuriuo naudojantis modeliui pateikiami tam tikri duomenys ir užduotis. Užduoties paskirtis yra atlikti kiekvienos operacijos semantinį konteksto ir elgsenos vertinimą pagal tam tikrus požymius bei

priskirti tam tikrus nurodytus kiekybinius rizikos įverčius. Šiai užduočiai atlikti yra pasitelkiamas GPT modelis naudojant API prieigą ir užklausų inžineriją formuluojant uždavinius.

Įprastoje finansinių mokėjimų aplinkoje dažniausiai taikytinos taisyklės akcentuoja įvairias skaitines reikšmes, pavyzdžiui, pinigų pervedimo dydis per tam tikrą laiko tarpą (dieną, savaitę, mėnesį, ketvirtį ar metus ir t.t.), mokėjimų dažnį per tuos pačius laiko intervalus ir t.t. Praktikoje, finansų įstaigos atliekant pinigų plovimo prevencijos procedūras kaip tekstinę informaciją vertina mokėjimų paskirties lauką, kuomet automatiškai tikrinama ar nėra juodajame sąrašė esančių raktažodžių, galinčių nurodyti įtartina ar nusikalstama veiklą. Taip pat vertinamos šalių geografinės lokacijos, atsižvelgiama iš kokios valstybės siunčiami pinigai, tačiau šis vertinimas yra paremtas vidine rizikos matrica, todėl iš esmės yra stipriai priklausomas nuo skaitinių reikšmių.

Kadangi modelis geba analizuoti tekstinę informaciją, tokią kaip veiklos kategorija, vietovės adresą, mokėjimo paskirtį, ir atpažinti semantinius ryšius, todėl ši informacija gali generuoti finansinių operacijų rizikos įverčius, galinčius padėti teisingai identifikuoti įtartinus mokėjimus. Naujausi moksliniai tyrimai nagrinėjantys didelių kalbos modelių (LLM) naudojimą sukčiavimo nustatymui srityse, kurios remiasi tekstine informacija (draudimo išmokos, finansinės ataskaitos, auditai) parodė reikšmingus rezultatus, rodančius, kad ši technologija ypač gerai gali suprasti tekstinę informaciją ir ją teisingai interpretuoti (Korkanti, 2024). Didelių kalbos modelių gebėjimas semantiškai nagrinėti tekstus, suprasti informaciją bei ją teisingai interpretuoti yra gana reikšmingas technologinis proveržis, galintis efektyviai papildyti sukčiavimo ir pinigų plovimo prevencijos procedūras. Dėl šių priežasčių pateiktame metode būtent siūloma praplėsti analizės imtį ir analizuoti ne tik skaitinę, bet ir tekstinę informaciją, siekiant nustatyti įtartina elgesį finansinių mokėjimų kontekste.

Siekiant tinkamai ir efektyviai pritaikyti didelį kalbos modelį (LLM) siūlomame metode, būtina suprasti įvesties/užklausos inžinerijos (angl. Prompt engineering) reikšmingumą sudarant užduotį modeliui, kad tekstinės informacijos interpretacija bei generuojami atsakymai būtų teisingi ir atitiktų analizuojamo konteksto tematiką. Pagal apibrėžimą, užklausa laikoma tekstu paremtą įvestį, kuri yra pateikiama modeliui, siekiant suteikti tam tikras gaires, pagal kurias modelis generuoja atsakymą (Marvin, Hellen, Jjingo ir Nakatumba-Nabende, 2024). Pagal 8 paveikslą, užklausų inžinerijos dalis užima svarbų vaidmenį didelio kalbos modelio atsakymų generavimo kontekste, kurios metu pateikiamos instrukcijos ir kontekstas modeliui.



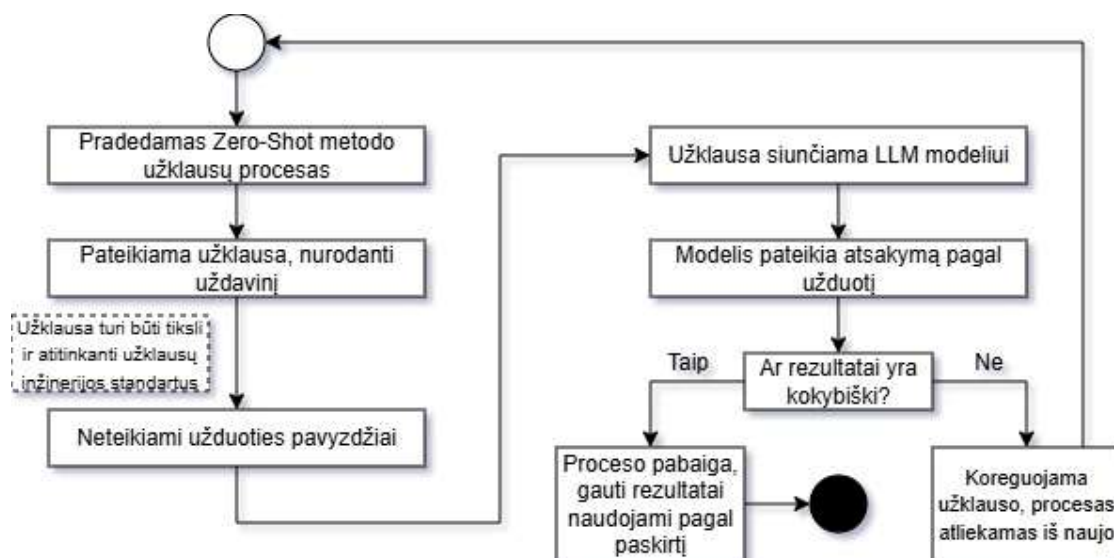
Šaltinis: Sahoo et. al. (2024). *A Systematic Survey of Prompt Engineering in Large Language Models: Techniques and Applications*, p. 1

pav. 9 Užklausų inžinerijos komponentų sąveikos schema

Anot Sahoo et. al. (2024), užklausų inžinerijos atsiradimas padarė labai reikšmingą įtaką didelių kalbų modelių transformatyvumui ir pritaikomumui įvairiose srityse. Naudojant šią techniką gauname galimybę, suteikiant instrukcijas ir kontekstą, keisti, modifikuoti didelio kalbos modelio generuojamą rezultatą ir jį koreguoti pagal instrukcijų pokyčius. Atsižvelgiant į šį kontekstą, būtent užklausų inžinerijos konceptas išplėtė tiek dirbtinio intelekto, tiek didelių kalbos modelių pritaikomumą. Ši technika leidžia turėti didesnę modelio kontrolę, todėl, galime plačiau eksperimentuoti ir su didesnio jautrumo duomenimis, pavyzdžiui, finansinius duomenis, ir gauti naujų bei naudingų įžvalgų, tuo pačiu išlaikant griežtą veiksmų kontrolę

Siūlomame metode naudojant didelio kalbos modelio (GPT) funkcionalumą suprantame, kad galutinio rezultato ir tyrimo kokybei didelę įtaką daro tinkamas užklausų inžinerijos taikymas. Suprasdami užklauskos kokybės svarbą galime užtikrinti ne tik kokybiškai sugeneruotą atsakymą, tačiau ir optimizuoti patiriamus kaštus bei tirti naujas technologijos taikymo būdus įvairiose srityse, tuo tarpu ir finansų sektoriuje (Marvin et. al., 2024). Tinkamam technologijos panaudojimui yra svarbu pasirinkti finansinių procesų analizei ir vertinimui tinkamą užklausų metodą. Pavyzdžiui, Sahoo et. al. (2024) nurodo net 12 skirtingų užklausų kategorijų, apimančių naujas užduotis be išsamaus mokymo (Zero-shot ir Few-shot užklausų metodai), mąstymo ir logikos kategoriją (Chain of Thought, Self-Consistency, Thread of Thought ir kiti metodai), haliucinacijų mažinimo kategorija ir kita. Platus užklausų metodikos pasirinkimas pagrindžia anksčiau minėtus teiginius, kad būtent užklausų inžinerija vaidina svarbų vaidmenį plačiame didelių kalbos modelių pritaikomume, sprendžiant įvairaus tipo, tame tarpe ir pinigų plovimo prevencijos problematiką.

Siekiant tinkamai ir efektyviai siūlomame metode pritaikyti didelio kalbos modelio technologinį funkcionalumą, kuris padėtų efektyviai nustatyti su sukčiavimu ir pinigų plovimu susijusius mokėjimų modelius (angl. Patterns) ir elgsenos bruožus, privalome pasirinkti tinkamą užklausų metodą. Atsižvelgiant į darbe keliamą problemą, tyrimo apimtį, duomenų rinkinius, maksimalius techninius kaštus ir naudojamą didelio kalbos modelio tipą (gpt-4.1-mini), galime teigti, kad metode vienas iš geriausiai pritaikomų užklausų metodų gali būti Zero-shot (liet. „nulinio-šūvio“) užklausų metodas. Autorius Li (2023) teigia, kad po GPT modelio sėkmės Zero-shot užklauskos tampa vis labiau patrauklesnės, kadangi kokybiškai pateikta užklausa be pavyzdžių gali padidinti modelio efektyvumą lyginant su Few-shot (liet. „kelių-šūvių“) ar kitais užklausų metodais. Būtent pavyzdžių pateikimo aspektas yra įvardijamas kaip esminis, kadangi dažnai modeliai pateiktus pavyzdžius interpretuoja ne kaip pavyzdį, o dalį pasakojamo naratyvo/konteksto, todėl galutiniai rezultatai gali reikšmingai skirtis. Visą Zero-shot užklauskos pateikimo darbo eigą galime apibūdinti pagal 9 paveiksle pateiktą schemą:

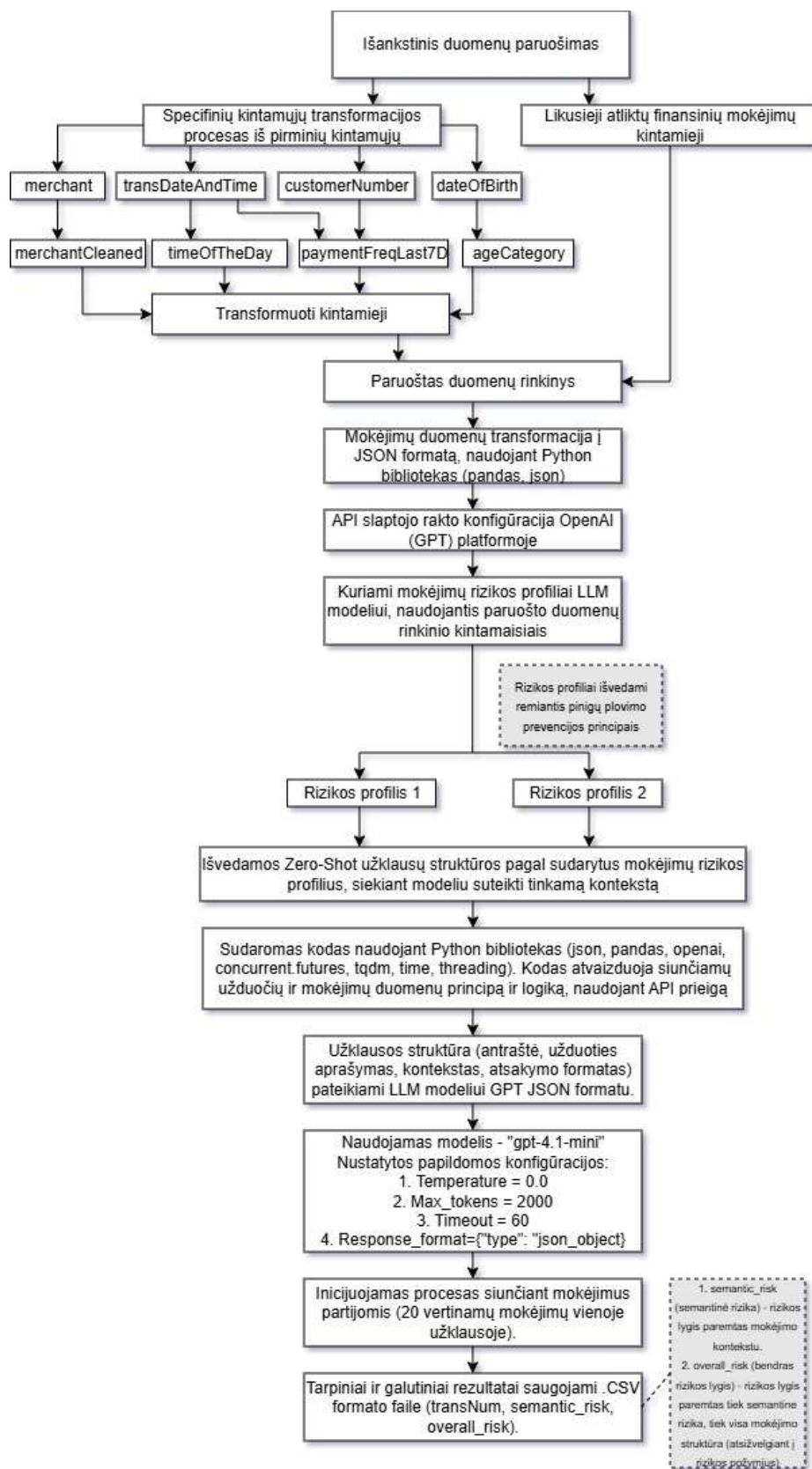


Šaltinis: sudaryta autoriaus

pav. 10 Zero-Shot užklauskos metodo veikimo schema

Zero-shot užklausų metodas susilaukia itin reikšmingo dėmesio, kadangi jie yra itin aiškūs, nereikalauja mokymo duomenų ir pavyzdžių, užklauskos projektavimo procesas yra paprastas, kadangi reikia tik suteikti aiškias instrukcijas su kontekstu ir pati užklauskos struktūra yra gana lanksti, ją galima koreguoti pagal gaunamus rezultatus (Li, 2023). Dėl šių priežasčių, įvertinus siūlomo metodo prieigą, analizuojamų finansinių duomenų turinį, nagrinėjamų finansinių įvykių kontekstą ir pasirinktą didelį kalbos modelį, galime teigti, kad Zero-shot užklauskos technika yra geriausiai tinkamas siūlomo metodo įgyvendinimui.

Išnagrinėjus esminius didelių kalbos modelių techninius aspektus ir remdamiesi tiek teorine, tiek praktine informacija atlikus mokslinių šaltinių analizę, siūlomo metodo didelio kalbos modelio funkcijai įgyvendinti bus naudojamas GPT modelis – **gpt-4.1-mini**. Siekdami sąveikauti su pasirinktu dideliu kalbos modeliu, naudosime GPT modelio aplikacijų programavimo sąsają (angl. Application programming interface, API). Metode siūloma pasirinkti aplikacijų programavimo sąsają, o ne įprastą vartotojo sąsają, kadangi API prieiga suteikia galimybę griežčiau valdyti modelio parametrus (temperatūrą, naudojamų žetonų kiekį), o tai leidžia valdyti pačio modelio elgseną, padidinti rezultatų nuspėjamumą ir kokybę. OpenAI siūlo klientų bibliotekas ir programinės įrangos kūrimo rinkinius (angl. Software development kits, SDKs), kurie yra pritaikyti visoms didžiausioms programavimo kalboms, todėl modelio API integracija į programinį kodą yra gana palengvinta. Atsižvelgiant į aukščiau pateiktas OpenAI siūlomas aplikacijų programavimo sąsajos techninius ypatumus ir kitus didelių kalbos modelių naudojimo aspektus (duomenų paruošimas, užklausos parengimas ir pan.), 11 paveiksle pateikiama išsami naudojamo didelio kalbos modelio – gpt-4.1-mini – veiksmų schema, kuri yra taikoma vertinant pateikiamų mokėjimų sukčiavimo rizikos kintamuosius ir išvedant naujus požymius (features) galinčius efektyviau nustatyti riziką, kad mokėjimas yra susijęs su sukčiavimu ar pinigų plovimo atvejais.



Šaltinis: sudaryta autoriaus

pav. 11 Siūlomo metodo didelio kalbos modelio veiksmų schema

Atliekant pradinio finansinių mokėjimų duomenų rinkinio išankstinį paruošimą, atliekame keturių kintamųjų: prekybininko (merchant), mokėjimo įvykio laiko žymos (transDateAndTime), kategorijos (category) ir kliento gimimo datos (dateOfBirth) transformacijas išvedant naujus kintamuosius. Naujas kintamasis merchantCleaned neturi „fraud“ žymės prie kiekvieno prekybininko pavadinimo, kadangi žyma neatitinka aplinkybių, kurios egzistuoja praktikoje bei darytų reikšmingą įtaką didelio kalbos modelio atliekamos tekstinės analizės proceso rezultatams. Kiti transformuoti kintamieji, tokie kaip dienos metas (timeOfTheDay), kliento atliekamų mokėjimų dažnis per praėjusią savaitę (paymentFreqLast7D) ir amžiaus kategorija (ageCategory) yra svarbūs sudarant mokėjimų rizikos profilius, kurie tekstone forma yra pateikiami kaip viena iš esminių konteksto informacijų atliekant pateiktą užduotį. Pasitelkiant Python programavimo kalbą ir naudingas bibliotekas, sukurtas programinis kodas veikdamas ciklo principu pateikia GPT modeliui užklausą su kontekstu, užduotimi ir atrinkta finansinių duomenų imtimi. Taip pat parašyto Python kodo struktūra taikome papildomas modelio (gpt-4.1-mini) konfigūracijas (modelio temperatūra, maksimalus simbolių kiekis, laiko limitas ir pateikiamų atsakymų formatas) kontroliuojame modelio veikimo specifikacijas, valdome laiko parametą, siekiant optimalios komunikacijos naudojant API bei kontroliuojame pateikiamų atsakymų formatą. Galiausiai, verta paminėti, kad modeliui įvertinti pateikiami mokėjimų duomenys rinkiniais, pavyzdžiui dvidešimt mokėjimų iš viso duomenų rinkinio. Taikant tokį veikimo principą išlaikome optimaliausią GPT modelio pajėgumą, vertinant finansinių duomenų rinkinį, kadangi visos apimties duomenų rinkinio įvedimas modeliui į konteksto parametrus gali lemti didesnę kiekį haliucinacijų ir kitų reikšmingų nukrypimų nuo užklauso.

Sekantis svarbus žingsnis siekiant išgauti aukštos kokybės rezultatus naudojant GPT modelį yra tinkamos struktūros užklausa Zero-shot metodui, pagal aukščiau aprašytas užklauso inžinerijos gaires. Kadangi siūlomas metodas yra orientuotas į finansinių operacijų stebėseną, siekiant teisingai identifikuoti sukčiavimo ir pinigų plovimo atvejus sumažinant klaidingai teigiamus įspėjimus, todėl užklauso turinys turi aiškiai apibrėžti finansinių duomenų ir situacijos kontekstą ir skirtą atlikti užduotį bei išvedamų rezultatų formą ir svarbą. Taip pat visą užklauso struktūrą formuojame pagal suformuotus mokėjimų rizikos profilius, kurie gali vertinti mokėjimų struktūras ir kontekstą iš skirtingų perspektyvų. Siūlomo metodo atveju siūloma pritaikyti du svarbius rizikos profilius pagal atliekamų mokėjimų ypatumus:

1. Kliento sociodemografinis rizikos profilis – sudaromas profilis pagal esminius kliento socialinius ir demografinius požymius, kuris paverčiamas tekstone struktūra didelio kalbos modelio analizei atlikti. Įvedant esminę kontekstinę informaciją naudojama ši rizikos profilio tekstinė struktūra: „Kliento profesija yra {job}, gimimo data {dateOfBirth}, amžiaus kategorija {ageCategory}, gyvenamoji vieta {city}, {state}. Finansinės operacijos suma yra {amount} USD, prekybininkas {merchantCleaned}, mokėjimo kategorija {category}“.

Naudojami pirminių ar transformuotų kintamųjų vietos laikiklius (angl. Placeholder) įdedama kiekvieno individualaus mokėjimo tiksli informacija, reikalinga kontekstui ir tolimesnei analizei.

2. Kliento finansinės elgsenos rizikos profilis – naudojami atrinkti finansinių mokėjimų kintamieji, kurie gali tiksliai apibūdinti kliento atliekamo mokėjimo ekonominį racionalumą. Šiam rizikos profiliui naudojama tekstinė struktūra: „*Mokėjimas, kurio dydis yra {amount} USD, atliktas {timeOfTheDay} metu, kuris buvo siunčiamas prekybininkui {merchantCleaned} už paslaugas {category}. Kliento per paskutinę parą atliktų mokėjimų skaičius buvo {paymentFrequencyLast7D}*“, o kliento amžiaus kategorija yra *{ageCategory}*“.

Remiantis užklausų inžinerijos gairėmis ir siekiant optimizuoti naudojamo GPT modelio darbą, žemiau pateikiama 2 lentelė ir 3 lentelė su tiksliomis užklausų struktūromis, kurios yra naudojamos eksperimento skyriuje vertinant finansinių mokėjimų duomenis ir atitinkamai įvertinant jų semantinę ir bendrinę riziką, susijusią su sukčiavimu:

2 lentelė

Sociodemografiniu kliento profilio šablonu paremta užklausos struktūra

Užklausos komponentas	Turinys
Užklausos antraštė/vaidmuo	Tu esi pinigų plovimo prevencijos ekspertas, kuris atlieka mokėjimų stebėseną ir vertina vykdomų mokėjimų riziką. Tavo tikslas yra analizuoti JAV fizinių klientų kortelinius mokėjimus, kurie juos vykdo tik JAV ribose.
Užduoties aprašymas	Atsižvelgiant į žemiau pateiktą sociodemografinį profilį, įvertink pateikiamos piniginės operacijos duomenis ir grąžink TIK validų JSON objektą su BŪTINAISIAIS laukais: 1. "semantic_risk": skaičius nuo 0.0 iki 1.0 (semantinė rizika pagal kontekstą - mokėjimo suderinamumas su kliento profiliu, amžiumi, profesija) 2. "overall_risk": skaičius nuo 0.0 iki 1.0 (bendras pinigų plovimo ir sukčiavimo balas, apimantis tiek semantinę, tiek bendrinę riziką). Įprastai turėtų būti didesnis nei "semantic_risk" Įverčiai gali varijuoti nuo 0.00 iki 1.00 (nuo 0.00 iki 0.25 - maža rizika, nuo 0.26 iki 0.65 - vidutinė rizika, nuo 0.66 iki 1.00 - didelė rizika) Vertinant sociodemografinį profilį ir mokėjimo duomenis, atkreipk dėmesį į mokėjimo riziką didinančius požymius: 1. Mokėjimo laikas - mokėjimas atliktas nakties metu (tarp 22:00 ir 05:00) didina riziką, ypač jeigu vykdomas už paslaugas, kurios dažniausiai perkamos kitu paros metu; 2. Kliento amžiaus kategorija (ageCategory) yra "Youth" arba "Senior; 3. Kliento amžius yra mažesnis nei 25 arba didesnis nei 65;

2 lentelės tęsinys

	4. Neįprastas sumos dydis kortelinių tipų mokėjimams, pavyzdžiui, vedama suma (amount) yra didesnė nei 500 JAV dolerių; 5. Neįprastai didelė suma atsižvelgiant į mokėjimo paskirtį, t.y. už ką atsiskaitoma (kinamasis category indikuoja mokėjimo paskirtį). Rizika kyla, jeigu sunku suprasti ekonominį racionalumą, pvz., už maisto prekes mokama 2000 dolerių;
Kontekstas	Kontekstas modeliui pateikiamas dviem segmentais: 1. Rizikos profilis – „Kliento profesija yra {job}, gimimo data {dateOfBirth}, amžiaus kategorija {ageCategory}, gyvenamoji vieta {city}, {state}. Finansinės operacijos suma yra {amount} USD, prekybininkas {merchantCleaned}, mokėjimo kategorija {category}.“ 2. Mokėjimų informacija, kuri yra pateikia dvidešimties mokėjimų rinkiniais JSON formatu. Pateikiami kintamieji: <i>transDateAndTime, customerNumber, merchantCleaned, category, amount, firstName, lastName, gender, street, city, state, zipCode, cityPopulation, job, dateOfBirth, transNum, timeOfDay, ageCategory, paymentFreqLast7D.</i>
Atsakymo formatas	Dabar tau pateikiamas KELIŲ operacijų sąrašas JSON formatu. Kiekvienai operacijai įvertink riziką ir grąžink rezultatą tokiu JSON formatu: <pre>"results": [{ "transNum": "<tas pats transNum kaip įvestyje>", "semantic_risk": <skaičius nuo 0.0 iki 1.0>, "overall_risk": <skaičius nuo 0.0 iki 1.0> },]</pre> SVARBU: 1. Grąžink TIK JSON objektą su vienu lauku "results". 2. "results" turi būti masyvas (list). Elementų seka turi atitikti įvesties transakcijų seką.

Šaltinis: sudaryta autoriaus

3 lentelė

Kliento finansinės elgsenos profilio šablonu paremta užklauso struktūra

Užklauso komponentas	Turinys
Užklauso antraštė/vaidmuo	Tu esi pinigų plovimo prevencijos ekspertas, kuris atlieka mokėjimų stebėseną ir vertina vykdomų mokėjimų riziką. Tavo tikslas yra analizuoti JAV fizinių klientų kortelinius mokėjimus, kurie juos vykdo tik JAV ribose.
Užduoties aprašymas	Atsižvelgiant į pateiktą finansinės elgsenos profilį, įvertink pateikiamos piniginių operacijos duomenis ir grąžink TIK validų JSON objektą su BŪTINAIŠ laukais:

	1. "semantic_risk": skaičius nuo 0.0 iki 1.0 (semantinė rizika pagal kontekstą - finansinio mokėjimo ekonominis racionalumas, suderinamumas su kita mokėjimo informacija) 2. "overall_risk": skaičius nuo 0.0 iki 1.0 (bendras pinigų plovimo ir sukčiavimo balas, apimantis tiek semantinę, tiek bendrinę riziką). Įprastai turėtų būti didesnis nei "semantic_risk" Įverčiai gali varijuoti nuo 0.00 iki 1.00 (nuo 0.00 iki 0.25 - maža rizika, nuo 0.26 iki 0.65 - vidutinė riziką, nuo 0.66 iki 1.00 - didelė rizika) Vertinant kliento finansinės elgsenos profilį ir mokėjimo duomenis, atkreipk dėmesį į mokėjimo riziką didinančius požymius: 1. Mokėjimas atliktas neįprastu paros metu (pavyzdžiui, už maisto prekes, buitines pirkinius, sveikatos priežiūros prekes nakties metu). 2. Neįprastai didelė operacijos suma pagal mokėjimo paskirtį (pavyzdžiui, už kurą ar maisto prekes sumokama 500 JAV dolerių ar daugiau). 3. Kliento amžiaus kategorija (ageCategory) yra "Youth" arba "Senior"; 4. Kliento amžius yra mažesnis nei 25 arba didesnis nei 65; 5. Didelis mokėjimų dažnis (75 ir daugiau) pagal paskutines 7 dienas, ypač atsižvelgiant į amžiaus kategoriją.
Kontekstas	Kontekstas modeliui pateikiamas dviem segmentais: 1. Rizikos profilis – <i>Mokėjimas, kurio dydis yra {amount} USD, atliktas {timeOfTheDay} metu, kuris buvo siunčiamas prekybininkui {merchantCleaned} už paslaugas {category}. Kliento per paskutinę parą atliktų mokėjimų skaičius buvo {paymentFrequencyLast7D}“, o kliento amžiaus kategorija yra {ageCategory}.</i> 2. Mokėjimų informacija, kuri yra pateikia dvidešimties mokėjimų rinkiniais JSON formatu. Pateikiami kintamieji: <i>transDateAndTime, customerNumber, merchantCleaned, category, amount, firstName, lastName, gender, street, city, state, zipCode, cityPopulation, job, dateOfBirth, transNum, timeOfTheDay, ageCategory, paymentFreqLast7D.</i>
Atsakymo formatas	Dabar tau pateikiamas KELIŲ operacijų sąrašas JSON formatu. Kiekvienai operacijai įvertink riziką ir grąžink rezultatą tokiu JSON formatu: <pre> "results": [{ "transNum": "<tas pats transNum kaip įvestyje>", "semantic_risk": <skaičius nuo 0.0 iki 1.0>, "overall_risk": <skaičius nuo 0.0 iki 1.0> },] </pre>

Šaltinis: sudaryta autoriaus

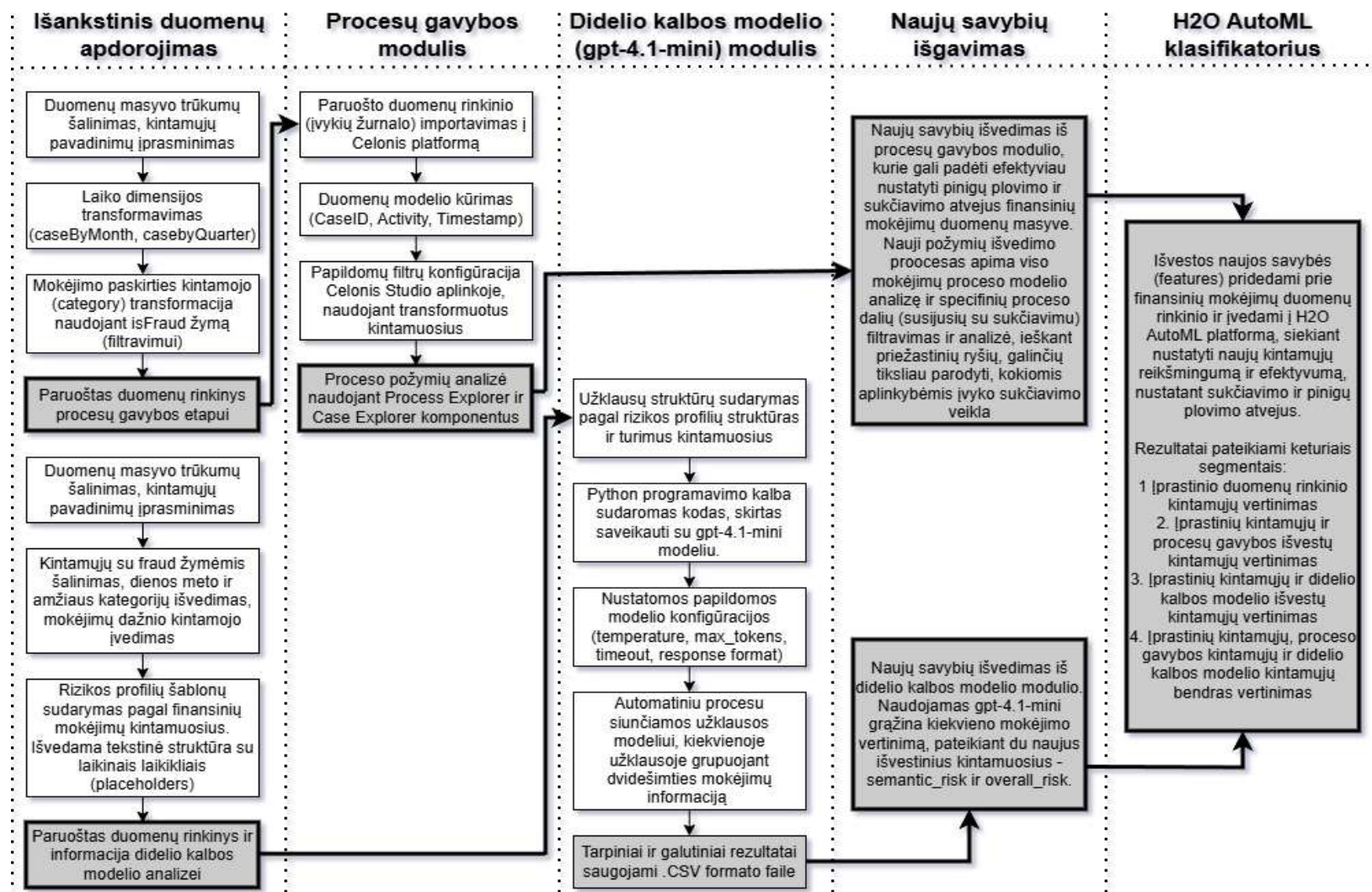
Naudodami pateiktas užklauso struktūras metodo vykdymo metu užtikriname, kad naudojamas gpt-4.1-mini modelis tinkamai supras kontekstą ir atitinkamai įvertins ir išves naujas mokėjimų savybes, t.y. papildomus vertinimus, kurie padėtų efektyviau ir tiksliau nustatyti įtartinas mokėjimų operacijas. Verta atkreipti dėmesį, kad modelio prašoma įvertinti kiekvieną mokėjimą pagal dvi savybes: semantinę riziką - pateiktos tekstinės informacijos (rizikos profilio ir kitų kintamųjų įvesties) vertinimą iš pinigų plovimo prevencijos perspektyvos ir bendrinę riziką, kuri išveda rizikos lygį, remiantis tiek semantiniu kontekstu, tiek kitais skaitiniais kintamaisiais, todėl vertinama iš bendrinės perspektyvos. Taip pat papildomam užduoties konteksto supratimui ir siekiant gauti tikslesnius rezultatus, išvedami papildomi požymiai (pagal sudarytus rizikos profilius), kurie gali didinti riziką, susijusią su pinigų plovimo ir sukčiavimo atvejais. Galiausiai, būtina pabrėžti, kad siekiant tinkamai kontroliuoti modelio išvedamų rezultatų kokybę ir stabilumą, koreguojame tam tikrus modelio parametrus, kurie yra nurodomi 11 paveiksle:

1. Temperature (liet. temperatūra) = 0.0 – atsitiktinumo ir kūrybiškumo parametras nustatomas į nulinę poziciją, siekiant gauti nuoseklius atsakymus ir maksimaliai sumažinti klaidų kiekį (pvz., sumažinamas šansas gauti netaisyklingos JSON struktūros atsakymą).
2. Max Tokens (liet. maksimalus žetonų (simbolių) kiekis) = 2 000 – skirta apriboti didelio kalbos modelio išvesties dydį, įvertinant gaunamų atsakymų dydžius, kaštus ir pan.
3. Timeout (liet. laiko limitas) = 60 – maksimalaus automatinio proceso laukimo laikas, laukiant atsakymo iš modelio prieš nutraukiant užklausą.
4. Response format (liet. atsakymo formatas) = {"type": "json_object"} – nurodymas modelio API, kad atsakymas privalo būti teisingas JSON objektas, o ne tekstas. Tai yra reikalinga siekiant efektyviai išlaikyti gaunamų atsakymų formato vientisumą ir neleidžia modeliui nukrypti nuo instrukcijų, taikant intensyvų ir laikui imlų procesą.

Visi nurodyti parametrų nustatymai yra būtini, norint efektyviai apdoroti didelius ir kompleksinius finansinių duomenų rinkinius, todėl privaloma taisyklingai vykdyti procesą, valdant išvesties formatą, veikimo logiką ir optimizuojant kaštus.

2.2.4 Metodo veikimo schema ir principai

Detaliai aprašytą metodo struktūrą galime išskirstyti į keturias veikimo dalis – išankstinis duomenų apdorojimas, procesų gavybos (Fuzzy Miner) taikymas, didelio kalbos modelio (GPT) taikymas ir naujų savybių išgavimas, kurių dalį gauname iš procesų gavybos sukurto procesų modelio, o dalį gauname sugeneruotą didelio kalbos modelio. Apibendrintą metodo veikimo struktūrą ir duomenų bei rezultatų išgavimo tvarką galime tiksliai atvaizduoti pagal schemą, pateiktą 12 paveiksle:



Šaltinis: sudaryta autoriaus

pav. 12 Bendra siūlomo metodo taikymo schema

Atliekant duomenų išankstinį apdorojimą, išvedami abejiems procesams (procesų gavybai ir didelio kalbos modelio vertinimui) reikalingi papildomi kintamieji, o dirbtinio intelekto atvejui papildomai parengiame semantinius tekstus – rizikos profilius ir bendros mokėjimų informacijos pateikimą JSON formatu, kurie yra reikalingi interpretacijai ir vertinimui atlikti. Atlikus išsamią literatūros analizę, procesų gavybos ir didelių kalbos modelių ypatumus, pasirinkta hibridinis metodas, kuomet abu analizės komponentai veikia lygiagrečiai, tačiau tyrimo eigoje šių komponentų rezultatai nėra perduodami vienas kitam, o sudaroma bendras naujų savybių (angl. Features) sąrašas. Šio sąrašo pagrindu atliekama tolimesnė analizė, naudojant H2O AutoML atviro kodo automatizuotą mašininio mokymosi sistemą, siekiant įvertinti naujai sukurtų ir sugeneruotų požymių efektyvumą ir aktualumą nustatant įtartinus mokėjimus visame duomenų rinkinyje.

Sudaryto metodo veikimo schema veikia hibridinio modelio principu, t.y. procesai nepersidengia, o vyksta lygiagrečiai, nes tiek procesų gavyba, tiek didelio kalbos modelio taikymas analizuoja skirtingus tų pačių duomenų elementus. Atlikus išsamią teorijos, literatūros ir kiekvieno proceso veikimo principus, įvairius taikymo būdus praktikoje bei jau esamus akademinius empirinių tyrimų rezultatus suprantame, kad procesų gavybos ir didelio kalbos modelio taikymas apima skirtingas duomenų imtis, skirtingus procesus ir leidžia atlikti skirtingo pobūdžio analizes. Procesų gavybos naudojamas algoritmas (Fuzzy miner) leidžia identifikuoti visą procesų modelį, o papildomi transformuoti kintamieji gali būti naudojami kaip filtravimo funkcija, kuri padeda analizuoti tiek įprastinę veiklą, tiek specifinius proceso srautus, kurie yra susiję su sukčiavimo operacijomis. Atlikta procesų ir atvejų analizė bei sukurti papildomi filtrai sukčiavimo procesui filtruoti, leidžia suformuoti naujas savybes, kurios yra pagrįstos identifikuotais požymiais (angl. Patterns), kurie yra dažniausiai pasitaikantys nusikalstamos veiklos kontekste.

Tuo tarpu naudojant didelį kalbos modelį (gpt-4.1-mini), galime atlikti papildomą tekstinės informacijos analizę, kuri nėra tokia reikšminga procesų gavybos kontekste. Būtent toks modelio panaudojimas leidžia atlikti mokėjimų stebėsenos analizę, į kurią įvedamas semantikos elementas, analizuojamas tekstas, vertinamas kontekstas, išvedami semantinės ir bendros rizikos kintamieji, atsižvelgiant į pateikiamus rizikos profilių, riziką didinančių požymių ir kitos mokėjimų informacijos komponentus. Tokia analizės priemonė nėra taikoma įprastinėje, taisyklėmis grįstose mokėjimų stebėsenos sistemose, todėl praplėsta analizuojamos informacijos apimtis gali leisti efektyviau aptikti su nusikalstama veikla susijusias operacijas. Dėl šių išvardintų priežasčių tiek procesų gavybos, tiek didelio kalbos modelio veikimas taikomas hibridiniu metodu, t.y. atliekama analizė ir jos rezultatai šiuose metodo komponentuose nepersidengia.

Galiausiai, būtina paminėti, kad gautų savybių sąrašas ir atitinkami papildomi duomenys yra įvedami atgal į pirminį finansinių mokėjimų duomenų rinkinį tolimesnei efektyvumo analizei. Siekiant metodiškai ir tiksliai nustatyti naujų požymių vertę (efektyvumą, tikslumą ir t.t.) analizuojant mokėjimų duomenis, tam tikslui naudojama H2O AutoML atviro kodo automatizuota mašininio mokymosi sistema, kuri atlieka automatizuotą modelių atranką, funkcijų inžineriją ir kuria prognozavimo modelis. Šį įrankį naudosime siekiant suprasti kaip efektyviai nauji požymiai koreliuoja su nustatytais sukčiavimo atvejais ir padeda juos teisingai identifikuoti. Atliktos efektyvumo analizės rezultatai išvedami keturiais skirtingais segmentais:

1. Įprastinių finansinių mokėjimų duomenų kintamųjų rinkinio analizė
2. Įprastinių finansinių mokėjimų kintamųjų rinkinys, papildytas išvestomis procesų gavybos savybėmis
3. Įprastinių finansinių mokėjimų kintamųjų rinkinys, papildytas gautomis savybėmis (mokėjimų įvertinimais) iš didelio kalbos modelio
4. Įprastinių finansinių mokėjimų kintamųjų rinkinys, papildytas tiek išvestomis procesų gavybos savybėmis, tiek vertinimais, gautais iš didelio kalbos modelio vertinimo

Šie segmentai leis atlikti detalią efektyvumo analizę, siekiant nustatyti kaip siūlomo metodo metu išgautos naujos savybės gali padidinti bendrą efektyvumą, nustatant sukčiavimo ir pinigų plovimo atvejus finansinių mokėjimų operacijose.

Taigi, apibendrinus, galime teigti, kad antrame skyriuje pasiūlytas metodas ir jo taikymo būdas gali būti, siekiant efektyviau nustatyti įtartinas mokėjimų operacijas finansinių duomenų rinkiniuose. Tiek teoriškai, tiek kitų mokslinių darbų empiriniais duomenimis pagrįstas hibridinis procesų gavybos algoritmo (Fuzzy Miner) ir didelio kalbos modelio (gpt-4.1-mini) metodas yra skirtas finansinių mokėjimų duomenų rinkinio analizei, o išvedamas rezultatas leidžia sudaryti papildomų savybių sąrašą, kurio įprastinės, taisyklėmis grįstos, stebėjimo sistemos negali sudaryti. Taip pat verta pabrėžti ir tai, kad siūlomas metodas apima tiek skaitinių, tiek tekstinių duomenų analizę, todėl gauti rezultatai apima didesnę kontekstą bei leidžia sudaryti papildomus požymius, skirtus įtartinoms operacijoms pažymėti.

Siūlomo metodo pradinėje stadijoje siūlomas duomenų paruošimo būdas – papildomų kintamųjų paruošimas maksimaliai sumažina duomenų triukšmo efektą, kas leidžia pasiekti efektyvesnį kitų komponentų darbą. Laiko žymų grupavimas į mėnesius ir ketvirčius, trūkstamų reikšmių šalinimas, nereikšmingų kintamųjų išvalymas bei papildomų kintamųjų su sukčiavimo žyme generavimas leidžia analizuoti globalesnę finansinių procesų eigos vaizdą, identifikuojant aiškius nukrypimus sugeneruotame modelyje. Šis žingsnis yra svarbus siekiant sumažinti „Spaghetti-type“ procesų modelio problemą. Taip pat atliktas tekstinės informacijos sisteminimas, tekstiniai rizikos profilio formavimai ir išsamus konteksto bei rizikos požymių įvedimas leidžia efektyviau dideliam

kalbos modeliui nustatyti kontekstą ir mokėjimų sentimentą bei generuoti aukštos kokybės įverčius, nurodančius mokėjimų rizikingumą.

Procesų gavybos etapui pasirinkimas taikyti Fuzzy Miner algoritmą yra pagrįstas finansinių duomenų kompleksiskumu ir apimtimi, todėl teigiama, kad būtent šis algoritmas gali veikti efektyviausiai nagrinėjant tokios struktūros duomenų masyvus. Išvesti papildomi kintamieji yra naudojami kaip papildomi filtrai, kurie leidžia efektyviau analizuoti tiek bendrą proceso modelį, tiek specifinius proceso etapus, kurie yra susiję su sukčiavimo atvejais. Pasitelkiant Celonis aplinkos siūlomus analitinius įrankius – Process Explorer ir Case Explorer – galime atlikti išsamią proceso ir atskirų atvejų analizes, o šis funkcionalumas yra svarbus, siekiant rasti priežastingumo ryšius, kurie dažniausiai lemia sukčiavimo ir pinigų plovimo įvykius.

Taip pat hibridiniu metodu, šalia procesų gavybos, taikoma didelio kalbos modelio analizė leidžia atlikti papildomą analizę naudojant tekstinę informaciją. Naudojant GPT modelį (gpt-4.1-mini), taikant Zero-shot principą, atliekamas visų mokėjimo operacijų semantinis vertinimas, kurio metu sukuriami kiti papildomi požymiai (įverčiai), galintys padidinti teisingai nustatomų sukčiavimo atvejų kiekį dėl platesnės analizės imties. Naudojant API prieigą ir taikant užklausų inžinerijos metodiką užtikrinama, kad modelis veiktų efektyviai, o generuojami požymiai atlieptų esamą kontekstą. Ši siūlomo metodo dalis suteikia papildomą interpretacinę dimensiją, kurios klasikiniai statistiniai ir taisyklėmis paremti metodai negali suteikti.

3. EKSPERIMENTINIS SKYRIUS

Šioje darbo dalyje atliekamas eksperimentas pagal pasiūlytą mokėjimų stebėsenos metodą, naudojant procesų gavybos neapibrėžto algoritmo (Fuzzy Miner) ir didelio kalbos modelio hibridinį modelį. Taip pat eksperimentiniame skyriuje aprašoma visa metodo proceso eiga remiantis pasiūlyto metodo skyriuje pateiktomis metodo veikimo schemomis. Galiausiai, vertinami eksperimento metu gauti rezultatai taikant klasifikatoriaus analizės metodą (angl. Cluster Analysis).

3.1 Siūlomo hibridinio metodo vykdymas

Šiame poskyryje atliekamas siūlomo hibridinio procesų gavybos ir didelio kalbos modelio siūlomas metodas, kurio metu siekiama išgauti naujų savybių kintamuosius, kurie gali padėti tiksliau identifikuoti sukčiavimo ir pinigų plovimo atvejus mokėjimų rinkinyje. Atliekant antrame skyriuje detaliai aprašyto metodo žingsnius, šiame poskyryje išsamiai aprašome atliekamus veiksmus paruošiant duomenų rinkinį analizei, Celonis procesų gavybos įrankių naudojimą, didelio kalbos modelio (gpt-4.1-mini) naudojimą per API sąsają ir gautų rezultatų – naujų savybių – išvedimą ir sisteminimą.

3.1.1. Duomenų rinkinys ir išankstinis duomenų paruošimas

Atliekant siūlomo hibridinio metodo eksperimentą, naudojame Sparkov Data Generation įrankio sugeneruotą atliktų finansinių mokėjimų duomenų rinkinį. Sparkov Data Generation įrankis taiko agentinio modeliavimo (angl. Agent-based modeling) principą, kuris imituoja individualių klientų ir prekybininkų sąveiką per realaus laiko perspektyvą. Šis įrankis naudoja tikimybinį elgesio generavimą ir kontroliuoja sukčiavimo atvejų įliejimo į įprastinių mokėjimų rinkinį kiekį. Taip pat įrankis yra paremtas realaus pasaulio mokėjimų ekosistemos principais, todėl generuojami mokėjimai yra glaudžiai susiję su įvairiais finansinės elgsenos aspektais. Dėl šių priežasčių reikšminga dalis akademikų savo moksliniuose darbuose naudoja šio įrankio sugeneruotus mokėjimų rinkinius. Verta atkreipti dėmesį, kad realūs finansinių mokėjimų duomenys yra itin jautri informacija, kurių valdymas ir saugojimas yra apibrėžtas įvairiuose reguliaciniuose reikalavimuose (pvz., Bendrasis duomenų apsaugos reglamentas, BDAR) ir standartuose (pvz., Mokėjimo kortelių pramonės duomenų saugumo standartas, PCI DSS). Dėl šių priežasčių siūlomo hibridinio metodo eksperimentui yra itin sunku gauti ir naudoti realiai įvykdytų asmenų/klientų mokėjimų, todėl remiamasi akademinius standartus atitinkančiais sintetinių duomenų rinkiniais. Verta paminėti, kad bendrinėje praktikoje, net ir realioje aplinkoje įvairios finansų institucijos vykdydamos įvairių finansinių produktų testavimus naudoja arba vidinius mokėjimų duomenų rinkinius, arba naudoja viešai prieinamus sintetinius finansinių mokėjimų rinkinius. Galiausiai, pažymime, kad šiame darbe gpt-

4.1-mini modelis pasirenkamas tik eksperimentiniams tikslams ir dėl turimų techninių išteklių ribotumo, todėl nesiūloma taikant šį metodą siųsti jautrius finansinių mokėjimų duomenis trečiosioms šalims (šiuo atveju OpenAI). Norint išvengti BDAR pažeidimų, realiomis sąlygomis finansų įstaiga turėtų naudoti vidinį ir lokaliai veikiantį didelį kalbos modelį arba pateikti mokėjimų informaciją trečiosioms šalims, neatskleidžiant duomenų, galinčių padėti identifikuoti mokėtoją ar gavėją.

Atliekant siūlomo hibridinio metodo eksperimentą naudojame Sparkov įrankio sugeneruotą duomenų rinkinį, kuris susideda iš 555 719 skirtingų mokėjimų, kuriuos atliko 924 skirtingi fiziniai asmenys 693-ims skirtingiems prekybininkams laikotarpiu nuo 2020 metų Birželio 11 dienos iki 2021 metų Sausio 1 dienos. Taip pat sugeneruotą finansinių mokėjimų duomenų rinkinį sudaro šie pirminiai kintamieji, kurie atvaizduoja pagrindinę su mokėjimo įvykiu susijusią informaciją: *Trans_date_trans_time; Cc_num; Merchant; Category; Amt; First; Last; Gender; Street; City; State; Zip; Lat; Long; City_pop; Job; Dob; Trans_num; Unix_time; Merch_lat; Merch_long; Is_Fraud*. Turimas duomenų rinkinys turi didžiąją dalį reikalingų savybių, kurios yra būtinos atlikti procesų gavybos ir didelio kalbos modelio analizės etapus, kadangi turi tiek laiko kintamąjį, tiek tekstinės informacijos kintamuosius. Nepaisant pirminio duomenų rinkinio tinkamumo eksperimentui atlikti, atliekame išankstinį duomenų paruošimo procesą, kuris apima tiek esamų kintamųjų koregavimą (kintamųjų pavadinimų įprasminimas, trūkstamų reikšmių valymas ir pan.), tiek naujų kintamųjų transformacijų, kurie specifiskai naudojami procesų gavybos ar didelio kalbos modelio vertinimo etapuose. Siekiant tinkamai atlikti esamų kintamųjų transformaciją, turime tiksliai išanalizuoti ir susisteminti turimų kintamųjų pavadinimus, įvertinti trūkstamų reikšmių kiekį, suprasti duomenų prasmingumą, todėl žemiau pateikiama 4 lentelė, įvardijanti kintamųjų pavadinimų pokyčius (siekiant įprasminti juos) ir jų suteikiamas reikšmes.

4 lentelė

Pirminių kintamųjų pavadinimų transformacijos ir reikšmės

Pirminis kintamasis	Galutinis kintamasis	Kintamojo reikšmė
<i>Trans_date_trans_time</i>	transDateAndTime	Reprezentuoja laiko žymą/įvykį, kuomet buvo atliktas mokėjimas. Svarbus procesų gavybos elementas
<i>Cc_num</i>	customerNumber	Unikalus kliento numeris. Svarbus kintamasis, siekiant atkurti mokėjimų veiksmų procesus.
<i>Merchant</i>	Merchant	Prekybininko pavadinimas. Tekstinės formos laukas, kurį galima panaudoti kontekstui.
<i>Category</i>	Category	Mokėjimo paskirtis, reprezentuojanti konkrečią veiklą/paskirtį, už kurią moka klientas. Svarbus kintamasis procesų gavybai ir didelio kalbos modelio analizei.

4 lentelės tęsinys

<i>Amt</i>	Amount	Atlikto mokėjimo suma. Kintamasis gali būti panaudotas kaip papildomas kontekstas ar išvedant kitus kintamuosius.
<i>First</i>	firstName	Kliento vardas (kontekstinė informacija)
<i>Last</i>	lastName	Kliento pavardė (kontekstinė informacija)
<i>Gender</i>	Gender	Kliento lytis
<i>Street</i>	Street	Kliento gyvenamosios vietos gatvės pavadinimas (demografinis kontekstas).
<i>City</i>	City	Kliento gyvenamosios vietos miesto pavadinimas (demografinis kontekstas).
<i>State</i>	State	Kliento gyvenamosios vietos valstijos pavadinimas (demografinis kontekstas).
<i>Zip</i>	zipCode	Kliento gyvenamosios vietos pašto kodas (demografinis kontekstas).
<i>Lat</i>	customerLatitude	Kliento geografinė platumą (demografinis kontekstas)
<i>Long</i>	customerLongitude	Kliento geografinė ilgumą (demografinis kontekstas)
<i>City_pop</i>	cityPopulation	Miesto (kuriame gyvena klientas) populiacija
<i>Job</i>	Job	Kliento oficialiai einamos pareigos (tinkamas kintamasis formuoti socialinį profilį)
<i>Dob</i>	dateOfBirth	Kliento gimimo data (socialinis kontekstas, papildomų kintamųjų išvedimui tinkama informacija)
<i>Trans_num</i>	transNum	Mokėjimo unikalūs numeris. Svarbus kintamasis siekiant teisingai susieti gautus procesų gavybos ir didelio kalbos modelio rezultatus (naujas savybes)
<i>Unix_time</i>	unixTime	Laiko atvaizdavimas sveikuoju skaičiumi (papildoma kontekstinė informacija)
<i>Merch_lat</i>	merchantLatitude	Prekybininko geografinė platumą (demografinis kontekstas)
<i>Merch_long</i>	merchantLongitude	Prekybininko geografinė ilgumą (demografinis kontekstas)
<i>Is_Fraud</i>	isFraud	Sukčiavimo žyma (atvaizduojama binarinėmis reikšmėmis 1 arba 0). Svarbus kintamasis, naudojant tiek eksperimento eigoje (procesų gavybos filtras), tiek galutinių rezultatų vertinime (H2O AutoML)

Šaltinis: sudaryta autoriaus.

Naudojantis Tableau Prep įrankiu atliktas esamų kintamųjų pavadinimų susistemimas ir įprasminimas bei atliktos pirminės analizės metu buvo nustatyta, kad finansinių mokėjimų duomenų

rinkinys neturi trūkstančių reikšmių, pvz., Null. Įvertinę turimų kintamųjų reikšmingumą ir naudą atliekamam eksperimentui, naudojantis apskaičiuoto lauko (angl. Calculated field) įrankiu, atliekame papildomų kintamųjų transformacijas, kurios yra reikalingos tiek procesų gavybos, tiek didelio kalbos modelio etapams. Žemiau pateikiama 5 lentelė, kurioje pateikiamos atliktos kintamųjų transformacijos:

Naujų papildomų kintamųjų transformacijos

Naudojami pirminiai kintamieji	Transformuotas kintamasis	Transformacija išreikšta programiniu kodu
customerNumber + transDateAndTime	caseByMonth	STR([customerNumber]) + '_' + STR(DATEPART('month', [transDateAndTime]))
customerNumber + transDateAndTime	caseByQuarter	STR([customerNumber]) + '_Q' + STR(DATEPART('quarter', [transDateAndTime]))
Category + isFraud	transformedCategory	IF([isFraud]=1) THEN [category]+"_fraud" ELSE [category] END
merchant	merchantCleaned	REGEXP_REPLACE([merchant], "^fraud_", "")
transDateAndTime	timeOfDay	IF DATEPART('hour', [transDateAndTime]) >= 22 OR DATEPART('hour', [transDateAndTime]) < 6 THEN "Night" ELSEIF DATEPART('hour', [transDateAndTime]) >= 6 AND DATEPART('hour', [transDateAndTime]) < 12 THEN "Morning" ELSEIF DATEPART('hour', [transDateAndTime]) >= 12 AND DATEPART('hour', [transDateAndTime]) < 18 THEN "Afternoon" ELSE "Evening" END
customerNumber + transDateAndTime + transNum	paymentFeqLast7D	df['transDateAndTime'] = pd.to_datetime(df['transDateAndTime']) df = df.sort_values(['customerNumber', 'transDateAndTime']) df['paymentsLast7Days'] = (df.groupby('customerNumber') .rolling('7D', on='transDateAndTime') .transNum.count() .reset_index(level=0, drop=True))
dateOfBirth	ageCategory	CASE WHEN DATEDIFF('year', [dateOfBirth], TODAY()) < 25 THEN "Youth" WHEN DATEDIFF('year', [dateOfBirth], TODAY()) >= 25

		AND DATEDIFF('year', [dateOfBirth], TODAY()) < 40 THEN "Young Adult" WHEN DATEDIFF('year', [dateOfBirth], TODAY()) >= 40 AND DATEDIFF('year', [dateOfBirth], TODAY()) < 65 THEN "Middle Age" WHEN DATEDIFF('year', [dateOfBirth], TODAY()) >= 65 THEN "Senior" END
--	--	---

Šaltinis: sudaryta autoriaus

5 lentelėje pateikiamos esamų kintamųjų transformacijos, kurių metu išvedami nauji kintamieji, naudojami procesų gavybos arba didelio kalbos modelio analizės kontekste. Šie išvesti kintamieji yra svarbūs dėl šių priežasčių:

1. **caseByMonth** – kintamasis leidžia kiekvieno kliento atliktus mokėjimus grupuoti vieno mėnesio intervalais. Toks grupavimas leidžia sudaryti aiškesnį ir mažiau komplikotą procesų modelį, kadangi be grupavimo kiekviena įvykio seka individualiai reprezentuoja vieną mokėjimą, todėl negalime atlikti išsamios proceso ir atvejų analizės.
2. **caseByQuarter** – kintamojo reikšmingumas yra toks pats kaip caseByMonth, tačiau šiuo atveju renkamės ketvirtinius intervalus, siekiant analizuoti kaip proceso modelis skiriasi per skirtingas laiko imtis.
3. **transformedCategory** – naudojant isFraud kintamąjį, išskiriame papildomas mokėjimo paskirtis su „fraud“ žyme. Tokia transformacija yra svarbi procesų gavybos taikyme, kuomet filtruojant tiek įprastinius, tiek sukčiavimo atvejus, galime matyti šių procesų pasiskirstymą, priežastingumo ryšius ir pan.
4. **merchantCleaned** – papildoma pirminio kintamojo transformacija, pašalinant „fraud“ žymes iš kiekvieno prekybininko pavadinimo. Ši žymė neturi jokio reikšmingo konteksto, todėl pašalinama, kadangi pateikiant prekybininko pavadinimus su „fraud“ žyme kaip kontekstą didelio kalbos modelio analizei, galime atsitiktinai daryti įtaką modelio vertinimui.
5. **timeOfDay** – išvestinis kategorinis kintamasis, kuris skirsto laiko žymę į dienos laiko kategoriją, pavyzdžiui, rytas, popietė, naktis ir pan. Formuojant tekstinius rizikos profilius didelio kalbos modelio analizei, toks kintamasis veikia kaip kintamasis, suteikiantis papildomą kontekstą, kadangi negalime užtikrinti, kad kiekvieno mokėjimo įvykdymo datą modelis teisingai interpretuos laiko kontekste.
6. **paymentFreqLast7D** – pagal kliento numerį (customerNumber), laiko žymą (transDateAndTime) ir unikalų mokėjimo numerį (transNum) išvedame naują kintamąjį, kuris

prie kiekvieno atlikto mokėjimo parodo, kiek klientas atliko mokėjimų per paskutines 7 dienas. Tai yra svarbus kintamasis, papildantis finansinės elgsenos rizikos profilį, todėl suteikiamas papildomas kontekstas didelio kalbos modelio analizei.

7. **ageCategory** – išvestinis kategorinis kintamasis, kuris sudaro klientų amžiaus kategoriją. Kintamojo paskirtis yra tokia pati kaip timeOfDay kintamojo, t.y. papildomo konteksto suteikimas formuojant tekstinį rizikos profilį, kadangi pateikus įprastinę gimimo datą, kyla rizika, kad didelis kalbos modelis neteisingai interpretuos žmogaus amžiaus grupę. Toks išvestinis kintamasis padeda sumažinti modelio haliucinacijų riziką.

Taip pat verta paminėti, kad beveik visos naujų kintamųjų transformacijos yra atliekamos naudojantis Tableau Prep įrankį, skirtą didelių duomenų valdymui, valymui ir transformacijai. Tačiau, dėl funkcinio ribotumo, paymentFreqLast7D kintamasis yra išvedamas naudojantis Python programavimo kalbos biblioteką „pandas“. Atlikus duomenų išankstinio paruošimo ir transformacijų etapą, atliekame sekančius siūlomo metodo etapus.

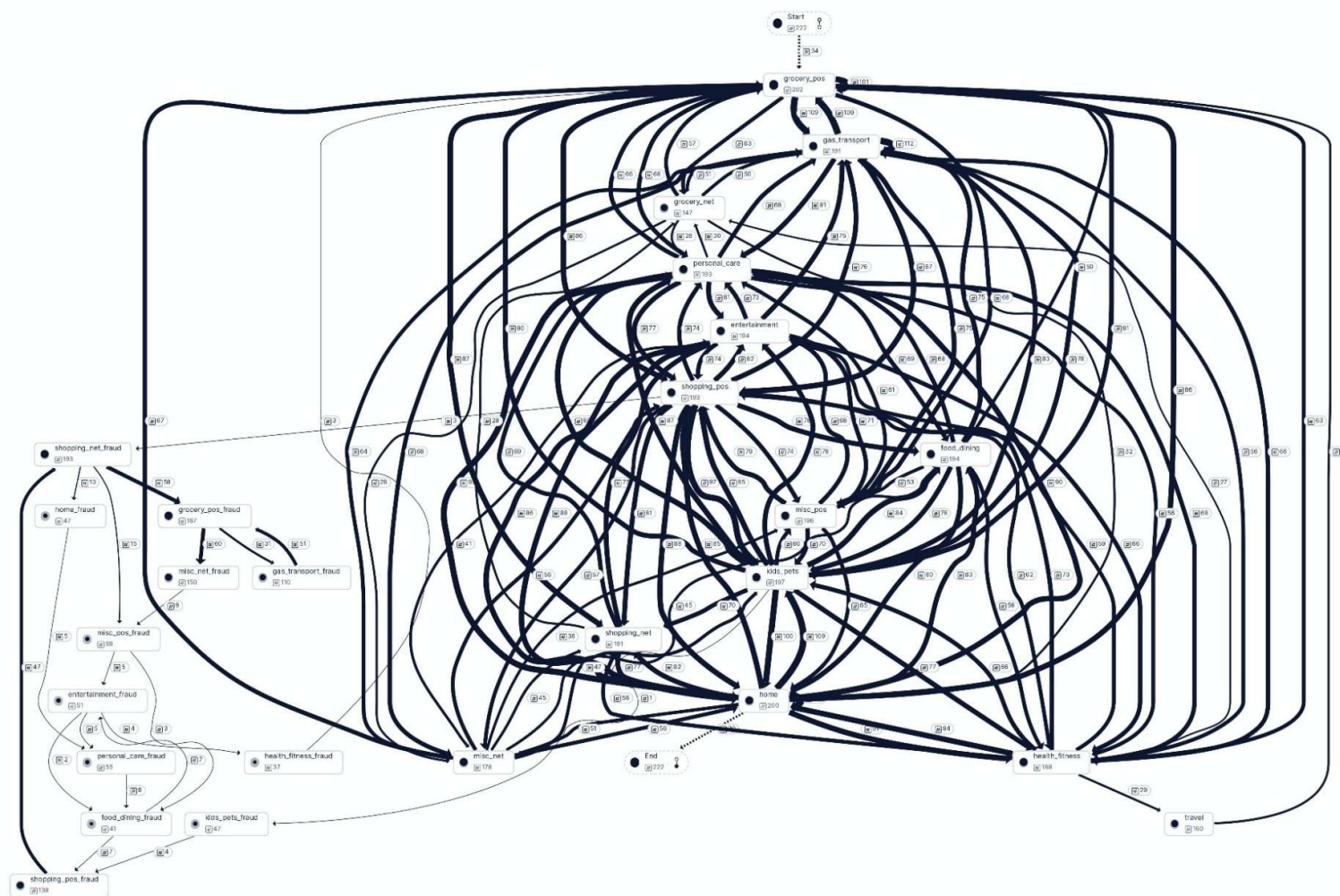
3.1.2. Procesų gavybos modulio taikymas Celonis platformoje

Atlikę išankstinį duomenų paruošimą ir reikiamas kintamųjų transformacijas į naujus kintamuosius, atliekame sekantį siūlomo metodo etapą – procesų gavybos taikymą Celonis platformoje. Kaip minėjome ankstesniame skyriuje, Celonis yra komercinė platforma, kuri tuo pačiu siūlo ir viešą prieigą prie procesų gavybos įrankio, kurio veikimo principas yra paremtas vidiniu procesų atradimo varikliu. Šis Celonis vidinio variklio veikimo principas yra paremtas neapibrėžto algoritmo (angl. Fuzzy Miner) veikimo principu, todėl įvertinus šio įrankio veikimo principą ir siūlomą funkcionalumą, jį naudojame finansinių mokėjimų procesų gavybos etapo vykdymui.

Verta paminėti, kad procesų gavybos etape naudojami ne visi anksčiau pateikti transformuoti kintamieji. Dėl šios priežasties toliau pateikiamas tikslus kintamųjų sąrašas, kurį naudojame procesų gavybos etapui Celonis platformoje: *caseByMonth* (antru atveju – *caseByQuarter*), *transDateAndTime*, *customerNumber*, *merchantCleaned*, *transformedCategory*, *amount*, *firstName*, *lastName*, *gender*, *street*, *city*, *state*, *zipCode*, *customerLatitude*, *customerLongitude*, *cityPopulation*, *job*, *dateOfBirth*, *transNum*, *unixTime*, *merchantLatitude*, *merchantLongitude*, *isFraud*. Celonis platformoje į susikurtą naują duomenų bazę (angl. New data pool) importuojame finansinių mokėjimų duomenų rinkinį sudarytą iš anksčiau paminėtų kintamųjų ir inicijuojame duomenų modelio (angl. Data model) generavimo etapą. Vykdam šį etapą svarbu teisingai pasirinkti tris esminius procesų modelio komponentus – atvejo identifikatorių (angl. Case ID), veiklą (angl. Activity) ir laiko žymą (angl. Timestamp). Atitinkamai formuodami procesų modelį, iš pateiktų kintamųjų atvejo identifikatoriumi pasirenkame *caseByMonth* (kitu atveju – *caseByQuarter*), veiklos komponentu renkame *transformedCategory*, o laiko žymą – *transDateAndTime*. Pasirinkę minėtus komponentus

gauname atvejo orientuotą (angl. Case-centric) procesų gavybos modelį. Naudojant *caseByMonth* ir *caseByQuarter* kintamuosius, finansiniai įvykiai yra grupuojami pagal atvejus, suskirstytus į vieno mėnesio ar ketvirčio intervalus, o tai leidžia generuoti du skirtingus procesų modelius pagal skirtingas laiko imtis. Toks pasirinkimas leidžia sukurti skirtingus procesų modelius, kurie pagal skirtingus laiko intervalus gali parodyti įvairius sukčiavimo atlikimo metodus, išskirti papildomas savybes, kurios yra imlios laiko perspektyvai. Tam tikri sukčiavimo modeliai trumpoje ar ilgoje laiko perspektyvoje gali būti sunkiau identifikuojami, todėl pasirenkama skirtingų laiko intervalų imtis leidžia išvengti analizės ribotumo.

Atlikus proceso modelių generavimo etapą ir susikūrę paketus, galime atlikti detalių modelių vertinimą ir analizę, naudojantis Celonis Studio aplinka. Šioje aplinkoje galime tirti procesų modelius ir jų atskirus etapus, naudojantis siūloma komponentų palete, kurioje šiuo atveju naudojame du svarbiausius komponentus – procesų naršyklę (angl. Process explorer) ir atvejų naršyklę (angl. Case explorer). Taip pat naudodami filtravimo įrankius atliekame procesų filtravimą ir naudodami *isFraud* kintamąjį išsifiltruojame tik tuos atvejų rinkinius, kuriuose įvyko sukčiavimo įvykiai, taip siekdami sumažinti modelio kompleksumą ir pašalinti įprastinius įvykius, kurie nėra susiję su sukčiavimu. Verta paminėti, kad išfiltruoti atvejai yra 1 mėnesio ar 1 ketvirčio intervale, todėl po kiekvienu sugrupuotu atveju yra nuo kelių iki kelių šimtų skirtingų įvykių (tiek įprastinių, tiek sukčiavimo). Dėl šios priežasties tyrimo metu galime analizuoti kokie veiksniai ar procesai gali lemti sukčiavimo atvejų atsiradimą finansinių mokėjimų sraute, naudojant skirtingas laiko kategorines imtis. Žemiau pateikiamas 13 paveikslas, kuriame atvaizduotas sugeneruotas procesų modelis (naudojant vieno mėnesio intervalus – *caseByMonth*), kuriam buvo pritaikyti aukščiau minėti filtrai:



Šaltinis: sudaryta autoriaus naudojant Celonis įrankius

pav. 13 Proceso modelis (Process Explorer) sudarytas pagal 1 mėnesio intervalą

13 paveiksle pavaizduotas sugeneruotas proceso modelis, sudarytas naudojant vieno mėnesio intervalus, naudojant Celonis procesų naršyklės komponentą, kuriame atvaizduotas filtruotų atvejų įvykių žemėlapis. Šiame modelyje matome, kaip sukčiavimo etiketėmis pažymėti procesų dalis įsimašo į įprastinių finansinių mokėjimų įvykių dalį, tačiau Celonis naudojamas neapibrėžtas algoritmas (angl. Fuzzy miner) leidžia atskirti šią dalį nuo proceso visumos. Toks funkcionalumas leidžia atlikti detalią analizę, nustatant esminius priežastingumo ryšius, kurie yra susiję su sukčiavimo įvykiais ir nustatyti požymius, kurie gali padėti lengviau identifikuoti sukčiavimo ir pinigų plovimo atvejus.. Panašiu principu, naudojant Celonis Studio aplinkoje esantį atvejų naršyklės komponentą gauname atskirą langą su vieno mėnesio atvejais, juose įvykusius tiek įprastinius finansinius mokėjimus, tiek sukčiavimo atvejus. Toks langas aiškia chronologine seka išdėsto kiekvieno kliento vykdytus mokėjimus per vieną mėnesį, laiko žymas, išsamius informacinius laukus, kurie yra sudaromi pagal pateiktą finansinių mokėjimų duomenų rinkinį. Kadangi, kaip aukščiau minėta, Celonis platformoje vykdoma atvejo orientuotą procesų gavybos modelį, todėl negalime tam tikrų komponentų informacijos eksportuoti .csv ar .xlsx formatu dėl platformos technologinių ribotumų, todėl 14 paveiksle pateikiamas sudarytos atvejų naršyklės lango pavyzdys, kuris naudojamas identifikuoti reikšmingus požymius, kurie gali būti susiję su sukčiavimo atvejais.

Case Id	# of Activities	Throughput Time	First Activity	First Activity Timestamp	Last Activity	Last Activity Timestamp	Case details: 676369110710_9
676369110710_9	116	30 d	misc_net	9/01/20 06:03:35	misc_pos	9/30/20 22:10:36	Search
676281772837_10	84	30 d	misc_net	10/01/20 03:53:35	travel	10/31/20 13:23:52	Activities 116 items
6597248210463122_10	41	30 d	health_fitness	10/01/20 12:32:32	home	10/31/20 19:11:32	travel 9/21/20 13:23:53 +20d
6593939509150446_9	36	28 d	misc_net	9/01/20 04:41:39	shopping_net	9/28/20 23:19:43	shopping_net 9/21/20 22:15:56 +21d
60485593109_9	42	28 d	shopping_pos	9/01/20 09:21:11	shopping_net_fraud	9/28/20 23:03:00	grocery_net 9/22/20 01:04:58 +21d
60487002085_9	48	29 d	home	9/01/20 16:43:35	kids_pets	9/30/20 23:37:48	grocery_pos 9/22/20 03:01:44 +21d
601147812335392_9	43	29 d	misc_net	9/01/20 07:51:31	grocery_pos_fraud	9/29/20 23:23:29	kids_pets 9/22/20 14:48:06 +21d
586100864972_9	12	1 d	kids_pets_fraud	9/23/20 18:34:20	entertainment_fraud	9/24/20 23:52:49	misc_pos 9/22/20 18:34:43 +22d
581508176315_10	37	29 d	personal_care	10/01/20 14:20:02	entertainment	10/30/20 16:16:40	shopping_pos 9/22/20 03:58:48 +22d
573283817795_10	31	27 d	health_fitness	10/04/20 12:21:52	kids_pets	10/31/20 19:02:57	gas_transport_fraud 9/23/20 05:40:54 +22d
570273151375_9	30	28 d	grocery_pos	9/02/20 07:32:24	health_fitness	9/29/20 21:10:47	grocery_pos_fraud 9/23/20 05:55:45 +22d
5456712664803820_10	60	29 d	home	10/01/20 23:28:11	gas_transport	10/31/20 09:49:43	misc_pos_fraud 9/23/20 19:28:16 +23d
4958589671582726883_10	94	30 d	shopping_net	10/01/20 00:54:48	personal_care	10/31/20 12:21:33	misc_net_fraud 9/24/20 00:07:53 +23d
4904681492230012_9	176	30 d	grocery_pos	9/01/20 03:52:41	personal_care	9/30/20 23:15:12	grocery_pos_fraud 9/24/20 00:10:32 +23d
4883407061576_10	9	1 d	shopping_net_fraud	10/24/20 20:34:24	shopping_net_fraud	10/25/20 23:06:13	grocery_pos_fraud 9/24/20 03:01:59 +23d
4861310130652566408_10	144	31 d	grocery_pos	10/01/20 01:57:04	shopping_pos	10/31/20 19:28:55	
4755696071492_10	126	31 d	grocery_pos	10/01/20 02:33:33	misc_pos	10/31/20 21:57:09	
47407131199405984_10	51	29 d	entertainment	10/01/20 01:45:27	grocery_net	10/30/20 05:48:06	
4710826438164847414_10	93	30 d	misc_net	10/01/20 09:20:09	health_fitness	10/31/20 16:32:11	
4646845581490336108_9	140	29 d	kids_pets	9/01/20 12:35:28	shopping_pos	9/30/20 14:51:08	
4623560839669_9	102	29 d	grocery_pos	9/01/20 11:08:21	kids_pets	9/30/20 19:48:50	
4599285557366057_10	82	30 d	shopping_pos	10/01/20 13:59:22	entertainment	10/31/20 19:29:27	
4586260469584_9	113	30 d	gas_transport	9/01/20 06:09:54	kids_pets	9/30/20 19:46:57	
4570636521433188_9	45	28 d	kids_pets	9/01/20 13:33:10	misc_pos	9/29/20 04:02:42	

Šaltinis: sudaryta autoriaus naudojant Celonis įrankius

pav. 14 Procesų modelio atvaizdavimas ir analizė naudojant atvejo naršyklės (Case explorer) langą

Toliau, 2 priede pateikiamas sudarytas procesų gavybos modelis, skirstant atvejus į ketvirtinius intervalus, t.y. kiekvienas sudarytas atvejis (pagal kliento identifikacinį numerį) savyje turi trijų mėnesių mokėjimus. Formuojant šį proceso modelį atvejo identifikatoriumi pasirenkame caseByQuarter, veiklos komponentu renkamės transformedCateogry, o laiko žymą –

transDateAndTime. Suformavę proceso modelį taikant kitą laiko intervalą, toliau sudarome atvejų naršyklės langą, pateiktą 3 priede, kuris naudojamas papildomai procesų analizei, siekiant nustatyti reikšmingas įvykių savybes, kurios padėtų efektyviau nustatyti sukčiavimo ir pinigų plovimo atvejus. Skirtingų procesų modelių sudarymas remiantis įvairiomis laiko imtimis yra svarbus, siekiant atlikti kuo detalesnę analizę, kuomet skirtingi laiko intervalai gali parodyti įvairesnius mokėjimų modelius (angl. patterns), sukčiavimo atvejus, kurie trumpesniuose ar ilgesniuose laiko variantuose gali būti geriau įlieti į bendrą mokėjimų srautą. Tokia praktika iš laiko intervalų perspektyvos yra taikoma ir įprastinėse mokėjimų stebėsenose, ypač retrospektyvios analizės procese.

Sudarius procesų modelius (naudojant finansinių mokėjimų duomenų rinkinius) ir naudojant procesų ir atvejų naršyklės komponentus, galime efektyviai analizuoti ir identifikuoti esminius finansinės elgsenos požymius, kurie yra susiję su sukčiavimo mokėjimų įvykiais. Tokie požymiai leidžia išvesti naujas savybes, kuriomis remiantis galime vertinti sukčiavimo atvejų nustatymo efektyvumą tolimesniame eksperimento etape. Atsižvelgiant į 13 ir 14 paveiksluose bei 2 ir 3 prieduose pateiktus ketvirčio laiko intervalo procesų naršyklės ir atvejų naršyklės rezultatus, atliekame analizę ir išvedame šias naujas savybes:

6 lentelė

Proceso gavybos metu išvestos savybės

Išvesta savybė	Savybės aprašymas	Programinis kodas
fromNetToPos	Atlikus proceso gavybos modeliavimą ir analizę buvo pastebėta, kad reikšminga sukčiavimo atvejų dalis įvyko pasikeičiant mokėjimo tipui (atlikus internetinį mokėjimą, sekantis mokėjimas atliekamas fiziškai pardavimo vietoje).	1 PRIEDAS
fromPosToNet	Tokia pati mokėjimų elgsena kaip aukščiau aprašytoje savybėje, tačiau veiksmų seka yra atvirkštinė (po atsiskaitymo prekybos vietoje kitas mokėjimas atliekamas internetu).	1 PRIEDAS
repeatedCategoryLast1H	Atvejų naršyklė parodė, kad reikšminga sukčiavimų dalis pasireiškia pasikartojančiuose to pačio tipo mokėjimuose sąlyginai dideliu dažniu per trumpą laiko tarpą (iki 1 val.).	1 PRIEDAS
sevenPlusCategoriesLast24H	Taip pat atvejų naršyklės komponentas padėjo nustatyti, kad dalis sukčiavimo atvejų pasireiškia ne tik dideliu mokėjimų dažniu, tačiau ir didele mokėjimo paskirties kaita.	1 PRIEDAS
personalCareAfterMiscNet	Proceso gavybos modeliai parodė, kad reikšminga dalis sukčiavimo mokėjimų įvyksta, kuomet vartotojas įvykdo atsiskaitymą už asmeninę priežiūrą, o po to atliekamas bendrinio pobūdžio mokėjimas internetu (ir atvirkščiai).	1 PRIEDAS

miscNetAfterPersonalCare	Savybė išvesta pagal aukščiau aprašytą savybę atvirkštinio pobūdžiu, nes sudaryti proceso modeliai parodė, kad tolimesnė atvejų seka nuveda į prieš tai buvusią kategoriją (asmeninę priežiūrą).	1 PRIEDAS
gasTransportToGroceryPos	Proceso gavybos modelių analizė parodė, kad nemaža dalis sukčiavimo atvejų įvyksta kuomet vykdomi atsiskaitymai už kurą, po kurių seka atsiskaitymas už maisto prekes prekybos vietoje (arba atvirkščiai).	1 PRIEDAS
groceryPosToGasTransport	Savybė išvesta pagal aukščiau aprašytą savybę atvirkštinio pobūdžiu, kadangi procesų gavybos modeliai parodė, kad reikšminga sukčiavimo įvykių dalis iš maisto produktų pirkimo grįžta į prieš tai buvusią kategoriją – mokėjimą už kurą.	1 PRIEDAS
ShoppingNetToGroceryPos ToMiscNet	Proceso gavybos modelių analizė parodė, kad tam tikra dalis sukčiavimo mokėjimų srauto įvyksta kuomet atliekamas internetinis mokėjimas už prekes, po jo atliekamas atsiskaitymas už maisto prekes prekybos vietoje, o galiausiai vėl atliekamas bendrinio pobūdžio mokėjimas internetu.	1 PRIEDAS
ShoppingNetToGroceryPos ToGasTransport	Proceso gavybos modelių analizė parodė, kad tam tikra dalis sukčiavimo mokėjimų srauto įvyksta kuomet atliekamas internetinis mokėjimas už prekes, po jo atliekamas atsiskaitymas už maisto prekes prekybos vietoje, o galiausiai atliekamas mokėjimas už kurą.	1 PRIEDAS
isNightTime	Procesų naršyklės komponentas padėjo nustatyti, kad didelė dalis mokėjimų pasikartoja naktiniu paros metu (tarp 22:00 val. ir 05:00 val.).	1 PRIEDAS

Šaltinis: sudaryta autoriaus

Proceso gavybos modelių analizės metu naudojami procesų ir atvejų naršyklių komponentai padėjo identifikuoti naujas reikšmingas savybes, kurios gali padėti prognozuoti sukčiavimo atvejus analizuojant didelius finansinių mokėjimų duomenų rinkinius. 6 lentelėje pateiktos savybės buvo išvestos naudojant Python programavimo kalbos pandas biblioteką. Ši biblioteka buvo pasirinkta savybių (požymių) išvedimui, nes ji leidžia efektyviai padoroti didelės apimties struktūrizuotus duomenis, atlikti slenkančio lango (angl. rolling window) skaičiavimo operacijas. Dėl šių funkcionalumų naudojame būtent pandas biblioteką, atlikti savybių inžinerijai (angl. features engineering), siekiant kompensuoti funkcinis ribotumus Tableau Prep aplinkoje. Verta paminėti, kad vienintelė savybė isNightTime buvo išvesta naudojant Tableau Prep įrankį. Išvestos savybės yra pridamos į pirminį finansinių mokėjimų duomenų rinkinį, kurį importavome į Celonis platformą, su tikslu atlikti šių kintamųjų efektyvumo analizę tolimesniame eksperimento etape (naudojant H2O AutoML mašininio mokymosi įrankį).

3.1.3. Didelio kalbos modelio (gpt-4.1-mini) modulio taikymas naudojant API sąsają

Atlikę išankstinį duomenų paruošimą ir reikiamas kintamųjų transformacijas paraleliai šalia procesų gavybos etapo atliekame kiekvieno mokėjimo sukčiavimo rizikos vertinimą naudojant GPT didelį kalbos modelį – gpt-4.1-mini – per API sąsają. Rizikos vertinimui naudojamos specifinės užklausos, paruošti kliento rizikos profiliai ir kita turimų finansinių mokėjimų rinkinio informacija. Taip pat būtina paminėti, kad didelio kalbos modelio analizei pateikiamas finansinių mokėjimų duomenų kintamųjų sąrašas yra dalinai skirtingas, lyginant su duomenų rinkiniais, pateiktu procesų gavybos etapui, todėl toliau pateikiami visi kintamieji, kuriuos modelis gauna mokėjimo vertinimui atlikti: *transDateAndTime*, *customerNumber*, *merchantCleaned*, *transformedCategory*, *amount*, *firstName*, *lastName*, *gender*, *street*, *city*, *state*, *zipCode*, *cityPopulation*, *job*, *dateOfBirth*, *transNum*, *timeOfTheDay*, *paymentFreqLast7D*, *ageCategory*. Galiausiai, vėl paminėti, kad kintamieji tokie kaip *customerLatitude*, *customerLongitude*, *unixTime*, *merchantLatitude* ir *merchantLongitude* yra pašalinami atsižvelgiant į jų mažą reikšmingumą ir siekiant sumažinti patiriamus kaštus atliekant mokėjimų semantinės ir bendros rizikos vertinimą naudojant didelį kalbos modelį.

Pirminis šio metodo etapo procesas yra visų turimų finansinių mokėjimų duomenų transformacija į JSON teksto formatą, kuris yra reikalingas, siekiant modeliui lengviau pateikti mokėjimų, kuriuos reikia įvertinti, informaciją. 1 priede pateikiamas išsamus Python kodas, kuris buvo sudarytas naudojantis *pandas* ir *json* bibliotekomis. Toks duomenų formatas leidžia patogiau ir efektyviau perduoti reikalingus duomenis GPT modelio (gpt-4.1-mini) vertinimui, atsižvelgiant į aplinkybes, kad duomenys yra perduodami grupuojant mokėjimų informaciją į rinkinius, kur kiekviename iš jų yra 30 unikalių mokėjimų. Mokėjimų rinkinio dydis pasirinktas atsižvelgiant į naudojamų darbuotojų (angl. *workers*) kiekį ir naudojamo modelio gpt-4.1-mini nustatytą API maksimalų naudojamų žetonų (angl. *tokens*) limitą.

Toliau, iš OpenAI susigeneruojame slapta API sąsajos raktą, kuris yra naudojamas 2 priede pateiktame Python programiniame kode, siekiant efektyviai komunikuoti su pasirinktu GPT modeliu (gpt-4.1-mini). Programiniam kodui parašyti buvo naudojamos *json*, *pandas*, *openai*, *concurrent.futures*, *tqdm*, *time* ir *threading* bibliotekos. Parašyto programinio kodo struktūrą ir logiką galime aprašyti šiuo eiliškumu:

1. **Duomenų įkėlimas ir sistemos konfigūracija** – šiame etape susiejame kodą su modeliu per API raktą, pasirenkame tikslų modelį (gpt-4.1-min), automatizuotam procesui nustatome užklausų ir sunaudojamų žetonų kiekį pagal oficialią OpenAI puslapyje esančią informaciją. Taip pat, siekiant sumažinti kaštus, modelio vertinimui siunčiame ne kiekvieną mokėjimą atskirai, o juos grupuojant į grupes po 30 mokėjimų, todėl programiniame kode nustatome maksimalų rinkinio dydį (*BATCH_SIZE=30*) ir darbininkų kiekį (*MAX_WORKERS=7*),

kurie siųs užklausas modeliui. Tuo tarpu BASE_PROMPT atvaizduoja pagrindinę užklaustos dalį, kurioje pateikiama užklaustos antraštė, užduotis, papildomas kontekstas.

2. **Užklaustos sudarymo logika** – sudarome rizikos profilio šabloną pagal nustatytas reikšmes (PROFILE_TEMPLATE), o šablone esančius vietas laikiklius (angl. Placeholders) pakeičia informacija iš kiekvieno vertinamo mokėjimo. Galiausiai get_batch_prompt() funkcija sudėlioja vieną bendrą užklausą, skirtą visam mokėjimų rinkiniui, t.y. kiekvienam rinkinio mokėjimui sudaromas rizikos profilis ir pateikiama pati mokėjimo informacija kaip papildomas kontekstas.
3. **API sąsajos valdymas** – kadangi visi GPT siūlomi modeliai turi tam tikrus užklausių ir žetonų limitus, todėl programinio kodo logikoje taikomas funkcionalumas wait_for_rate_limit(), kuris leidžia apskaičiuoti paruošto užklaustos ir mokėjimų rinkinio bendrą žetonų sumą ir atitinkamai pagal limitus sprendžiama ar užklausa gali būti pateikta modeliui. Automatinis procesas yra stabdomas sleep() ir threading.Lock() funkcijomis, siekiant valdyti limitų viršijimą ir kelių procesų veikimą vienu metu.
4. **Mokėjimų rinkinio apdorojimo mechanizmas** – naudodami process_batch() funkciją, mes atliekame skaičiavimus, kiek tokenų sudaro visas rinkinys, ar jis atitinka modelio API limitus. Jeigu ši sąlyga yra patenkinama, tuomet visas rinkinys siunčiamas modeliui (gpt-4.1-mini) per OpenAI API sąsają su užduotimi grąžinti kiekvieno mokėjimo įvertinimą (semantic_risk ir overall_risk) JSON formatu. Minėta funkcija taip pat tikrina ar modelio pateikiamas atsakymas atitinka tikimasi duomenų formatą ir ar visiems mokėjimams visi įverčiai buvo grąžinti. Jeigu bent vienas mokėjimas negauna tinkamo ar pilno vertinimo, tuomet inicijuojama kita funkcija – call_single_txn() – kuri dar kartą išsiunčia pilną užklausa su vienu mokėjimu ir rezultatas yra grąžinamas atskirai. Taip pat ši funkcija suveikia ir tuomet, kai modelis atsakymą pateikia neteisingu eiliškumu, negrąžina viso rinkinio ar nutraukia atsakymą.
5. **Lygiagrečių procesų mechanizmas** – atsižvelgiant į tai, kad vertinamų mokėjimų duomenų apimtis yra gana didelė, todėl siekiant išvengti ilgo proceso ir didesnių nei numatyta kaštų, pasirenkame vykdyti kelis procesus vienu metu (ThreadPoolExecutor(max_workers=MAX_WORKERS)). Kiekvienam darbininkui yra suteikiami unikalūs mokėjimų rinkiniai, o siekiant suvaldyti užklausių procesą, kiekvienas darbininkas laukia leidimo, kuomet gali pateikti užklausą per API sąsają.
6. **Tarpinis rezultatų surinkimas** – kaip minėjome anksčiau, didelio kalbos modelio mokėjimų vertinimo procesas yra laikui imlus procesas, todėl siekiant išvengti nenumatytų klaidų ar proceso nutrūkimo, taikome papildomą tarpinių rezultatų saugojimo funkciją (CHECKPOINT_EVENT = 5000). T.y. kas penkis tūkstančius įvertintų mokėjimų programa

išsaugo tarpinius rezultatus (mokėjimo identifikacinį numerį, `semantic_risk`, `overall_risk`) į .csv formato failą.

- Galutinis rezultatų išvedimas** – galiausiai, kuomet paskutinis vertinamų mokėjimų rinkinys yra įvertintas ir jo rezultatai teisingai pateikiami, tuomet naudojant `pd.DataFrame()` funkciją išvedame pilną įvertintų mokėjimų rezultatų failą .csv formate, kuris yra naudojamas tolimesniame eksperimento etape (rezultatų vertinime).

Sukurtas programinis kodas leidžia automatiškai apdoroti masyvius įvykdytų mokėjimų duomenų rinkinius, o papildomai įvesti loginiai apribojimai veikia kaip reikšmingas kaštų optimizavimo įrankis, tiek iš laiko, tiek iš finansinės perspektyvos.

Antroje eksperimento dalyje atliekame mokėjimų semantinės ir bendrinės rizikos vertinimą, naudojant `gpt-4.1-mini` modelį. Naudojant 1 priede pateiktą nuorodą į Github platformą, kurioje patalpintas programinis kodas, buvo pateikti penki šimtai penkiasdešimt penki tūkstančiai ir septyni šimtai devyniolika (555 719) mokėjimų, o procesas buvo kartojamas du kartus, taikant anksčiau pateiktas dviejų užklausų struktūras. Kiekvienam iš įvertintų mokėjimų, naudojamas modelis `gpt-4.1-mini`, pateikė semantinės rizikos (`semantic_risk`) ir bendros rizikos (`overall_risk`) įverčius, režiuose nuo 0.00 iki 1.00. Pagal gautus rezultatus sistemiškai išvedame papildomas savybes, pagal `gpt-4.1-mini` modelio pateiktus įverčius kiekvienam mokėjimui, kurie yra pavaizduoti 7 lentelėje:

7 lentelė

Didelio kalbos modelio mokėjimų įverčių savybės

Savybė (pagal įverčius)	Savybės kontekstinis paaiškinimas
<code>semanticRiskRP1</code>	Semantinės rizikos įverčiai kiekvienam mokėjimui, kurie buvo įvertinti pagal 3 lentelėje pateiktą užklausos struktūrą ir kontekstą. Ši užklausos struktūra yra suformuluota pagal sociodemografinį kliento rizikos profilį.
<code>overallRiskRP1</code>	Bendrinės rizikos įverčiai kiekvienam mokėjimui, kurie buvo įvertinti pagal 3 lentelėje pateiktą užklausos struktūrą ir kontekstą, remiantis sociodemografiniu kliento rizikos profiliu.
<code>semanticRiskRP2</code>	Semantinės rizikos įverčiai kiekvienam mokėjimui, kurie buvo įvertinti pagal 4 lentelėje pateiktą užklausos struktūrą ir kontekstą. Ši užklausos struktūra buvo suformuluota pagal kliento finansinės elgsenos rizikos profilį.
<code>overallRiskRP2</code>	Bendrinės rizikos įverčiai kiekvienam mokėjimui, kurie buvo įvertinti pagal 4 lentelėje pateiktą užklausos struktūrą ir kontekstą, remiantis kliento finansinės elgsenos rizikos profiliu.

Šaltinis: sudaryta autoriaus

Šias reikšmes toliau naudojame kaip papildomas savybes, siekiant efektyviau nustatyti sukčiavimo ir pinigų plovimo atvejus mokėjimų duomenų sraute, kas leistų sumažinti sugeneruojamų klaidingai

teigiamų įspėjimų kiekį. Galiausiai, atliekant gautų rezultatų tolimesnę analizę Tableau Prep aplinkoje, nustatėme, jog modelis grąžino įverčius visiems mokėjimams, t.y. naudojant tinkamą užklauso struktūrą ir pasirinkus tinkamus modelio nustatymus visiškai išvengėme haliucinacijų problemos (jokių reikšmių už nustatytų rėžių, jokių Null reikšmių ir pan).

3.2 Rezultatų analizė naudojant H2O AutoML įrankį

Atlikę procesų gavybos analizę naudojant Celonis gavybos įrankį ir mokėjimų rizikingumo vertinimą naudojantis gpt-4.1-mini dideliu kalbos modeliu, išvedame naujas savybes, kurios gali padėti pagerinti sukčiavimo ir pinigų plovimo atvejų identifikavimą dideliuose mokėjimų srautuose, taip sumažinant klaidingai teigiamų įspėjimų kiekį. Šių naujų savybių daromos įtakos bendram efektyvumui vertinimą atliekame naudojanti H2O AutoML atviro kodo automatizuotą mašininio mokymosi sistemą.

Prieš atliekant efektyvumo analizę naudojant H2O AutoML įrankį, paruošiame keturis pagrindinius duomenų rinkinius, kurie yra naudojami analizės metu. Paruošti duomenų rinkiniai yra suskirstomi į keturias kategorijas:

1. Pradinis duomenų rinkinys.
2. Pradinis duomenų rinkinys, papildytas gautomis savybėmis iš proceso gavybos modelių.
3. Pradinis duomenų rinkinys, papildytas gautomis savybėmis (įverčiais) iš didelio kalbos modelio.
4. Pradinis duomenų rinkinys, papildytas tiek savybėmis, gautomis iš proceso gavybos etapo, tiek savybėmis (įverčiais), gautais iš didelio kalbos modelio.

Duomenų rinkinius išskirdami į keturias kategorijas galime įvertinti, kaip atskirų metodo procesų komponentai (procesų gavyba ir didelio kalbos modelio taikymas) daro įtaką efektyvumui, t.y. kaip papildomos išvestos savybės padeda algoritmui teisingai nustatyti sukčiavimo atvejus ir generuoti kuo mažiau klaidingai teigiamų perspėjimų.

Tolimesniame etape atliekame aukščiau paminėtų duomenų importavimą į H2O AutoML įrankį ir siekiant užtikrinti tinkamą algoritmo veikimą, kintamojo isFraud ir kitų išvestų savybių tipą pakeičiame į kategorinį tipą (angl. enum), kadangi visi šie kintamieji yra kategorinio tipo (kintamųjų tipo pakeitimo pavyzdys pateikiamas 4 priede). Toliau atliekame duomenų rinkinio dalybą, dalinant 70 proc. duomenų į mokymosi (angl. training) kategoriją, o likusius 30 proc. duomenų paliekame validavimo (angl. validation) etapui (žr. 5 priedą). Efektyvumo analizės etapui atliekame eksperimentus naudojant du modelius – Gradient Boosting Machine (GBM) ir Distributed Random Forest (DRF), kurie yra laikomi kaip sprendimo medžių pagrindu veikiantys modeliai. Kiti modeliai eksperimento etape nebuvo naudojami, siekiant išlaikyti galutinių rezultatų efektyvumo vertinimo

vientisumą, kadangi kiti modeliai turi skirtingus parametrus, todėl kyla rizika, kad galutiniai rezultatai bus palyginti neteisingai. Anksčiau minėti modeliai pasirenkami dėl jų tinkamumo analizuoti didelius duomenų rinkinius, kurių duomenys turi ne tiesioginius ryšius, yra triukšmingi ir matomas ryškus klasių disbalansas (šiuo atveju isFraud) bei visi išskirti algoritmai yra tinkama atlikti klasifikaciją. Siekiant atrasti optimaliausią kiekvieno algoritmo techninių parametrų kombinaciją, kuri padėtų pasiekti optimaliausius rezultatus, naudojantis H2O AutoML įrankiu kiekvienos iteracijos metu pateikiame vieną iš keturių sudarytų duomenų rinkinių ir mokymosi etape atliekame papildomas tris iteracijas, pritaikant skirtingas algoritmo parametrų kombinacijas. Pagrindiniai parametrai, kuriuos keičiame iteracijų metu yra šie: medžių kiekis (ntrees), maksimalus medžių gylis (max_depth), minimalus stebėjimų skaičius (min_rows), mokymosi greitis (learn_rate), klasių balanso stiprumas (class_sampling_factor), balansavimo dydis (max_after_balance_size). Gautus tarpinius rezultatus išvedame trijose atskirose lentelėse, pateikiant esminius naudojamus parametrus ir atitinkamai gautus modelių rodiklius.

Naudojant Distributed Random Forest modelį, jį apmokome naudojant keturis skirtingus duomenų rinkinius, taikant tris skirtingas parametrų variacijas, kurių metu keičiame medžių kiekį, gylį, klasių svorius, klasių balanso dydžius ir kitus parametrus. Toks būdas taikomas, siekiant nustatyti optimaliausius parametrus, kurie leistų modeliui pasiekti kuo efektyvesnį galutinį rezultatą, t.y. sugeneruoti kuo mažiau klaidingai teigiamų įspėjimų. Žemiau pateikiama 8 lentelė, kurioje pateikiami skirtingos modelio parametrų sąrankos ir gauti rezultatai:

Naudojamo DRF modelio pirminių rezultatų analizė taikant skirtingus parametrus

Duomenų rinkinys (DR)	Modelis	Medžiai	Gylis	Stebėjimų skaičius	Mokymosi greitis	Klasių Balansas	Balanso dydis	AUC	PR_AUC	R ²	False Positives	FP rate	F1 Score
Pradinis duomenų rinkinys	DRF	15	5	10	N/A	1, 1	5	0,98	0,45	0,07	373	49,00%	0,55
Pradinis duomenų rinkinys	DRF	35	10	15	N/A	1, 3	5	0,96	0,63	0,32	144	26,60%	0,67
Pradinis duomenų rinkinys	DRF	50	7	20	N/A	1, 5	5	0,97	0,63	0,38	125	24,30%	0,67
Pradinis DR su procesų gavybos savybėmis	DRF	15	5	10	N/A	1, 1	5	0,96	0,62	0,34	209	33,8%	0,63
Pradinis DR su procesų gavybos savybėmis	DRF	35	10	15	N/A	1, 3	5	0,97	0,67	0,37	106	20,50%	0,69
Pradinis DR su procesų gavybos savybėmis	DRF	50	7	20	N/A	1, 5	5	0,97	0,69	0,3	114	20,80%	0,71
Pradinis DR su didelio kalbos modelio savybėmis	DRF	15	5	10	N/A	1, 1	5	0,97	0,57	0,16	259	39,90%	0,59
Pradinis DR su didelio kalbos modelio savybėmis	DRF	35	10	15	N/A	1, 3	5	0,97	0,69	0,43	115	21,60%	0,7
Pradinis DR su didelio kalbos modelio savybėmis	DRF	50	7	20	N/A	1, 5	5	0,98	0,69	0,43	118	23%	0,67
Pradinis DR su visomis savybėmis	DRF	15	5	10	N/A	1, 1	5	0,98	0,56	0,13	185	32,50%	0,62
Pradinis DR su visomis savybėmis	DRF	35	10	15	N/A	1, 3	5	0,98	0,69	0,42	82	17,80%	0,67
Pradinis DR su visomis savybėmis	DRF	50	7	20	N/A	1, 5	5	0,98	0,73	0,43	98	18,50%	0,72

Šaltinis: sudaryta autoriaus

Pagal 8 lentelėje pateiktus pirminius rezultatus, vertiname kaip skirtingi parametrai veikia Distributed Random Forest (DRF) efektyvumą, generuojant mažesnę klaidingai teigiamų įspėjimų kiekį. Pagal gautus rezultatus matome, kad klaidingai teigiamų įspėjimų kiekiui didelę reikšmę turi tiek pasirenkamų medžių kiekis, tiek medžių gylis, minimalus stebėjimų skaičius mazge ir klasių balansavimas. Pavyzdžiui, taikant parametrų kombinaciją, kuri susideda iš 15 medžių, 5 mazgų gylio, 10 minimalių stebėjimų mazge bei nustatant balanso dydžio reikšmę kaip 5, visuose duomenų rinkiniuose gauname prasčiausius tiek klaidingai teigiamų įspėjimų kiekius (varijuoja tarp 373 ir 185), tiek procentinės klaidingai teigiamų įspėjimų dalies nuo visų sugeneruotų įspėjimų rezultatus. Taip pat ši parametrų kombinacija laikytina prasčiausia iš efektyvumo perspektyvos dėl itin prastų kitų modelio rodiklių, pavyzdžiui determinacijos koeficientas (R^2) varijavo vos tarp 0,07 ir 0,34, F1 reikšmė parodė prasčiausią tikslumo ir jautrumo balansą (įverčiai tarp 0,55 ir 0,63), o PR_AUC kreivės ploto rodiklis parodė, kad naudojant šiuos parametrus modelis prasčiausiai atpažįsta retą klasę, todėl generuoja daugiausiai klaidingai teigiamų įspėjimų.

Tuo tarpu, taikant parametrų kombinaciją, kuri susideda iš 35 medžių, 10 mazgų gylio, 15 minimalių stebėjimų mazge ir padidinant klasių svorius iki 1:3 reikšmingumo gauname žymiai geresnius rezultatus. Ryškiausiai šios parametrų kombinacijos efektyvumas atsispindi per sugeneruojamų klaidingai teigiamų įspėjimų kiekį, kurį modelis eksperimento metu fiksavo tarp 144 (naudojant tik pradinį duomenų rinkinį) ir 82 (naudojant pradinį duomenų rinkinį su savybėmis iš kitų metodo procesų). Taip pat naudojant šią parametrų kombinaciją matome ryškų klaidingai teigiamų įspėjimų procentinės dalies nuo visų sugeneruojamų įspėjimų sumažėjimą visuose duomenų rinkiniuose (matomas sumažėjimas nuo 13,3% iki 22,4%). Kiti modelio rodikliai indikuoja žymiai efektyvesnį modelio veikimą, mokantis iš pateiktų duomenų rinkinių, pavyzdžiui determinacijos koeficientas (R^2) varijuoja tarp 0,32 ir 0,43, PR_AUC kreivės ploto rodiklis varijuoja tarp 0,63 iki 0,69, o F1 reikšmė rodo, kad modelis žymiai geriau atpažįsta retą klasę (varijuoja tarp 0,67 ir 0,71), kas pagrindžia mažesnę sugeneruojamų klaidingai teigiamų įspėjimų kiekį. Paskutinė modelio parametrų kombinacija reikšmingo rodiklių pagerėjimo neparodė ir visuose, apart pirminio, duomenų rinkiniuose parodė didesnę sugeneruojamų klaidingai teigiamų įspėjimų kiekį (varijuoja tarp 125 ir 98), todėl pirminių rezultatų analizė leidžia teigti, kad antra modelio parametrų kombinacija leidžia Distributed Random Forest modeliui pateikti efektyviausius rezultatus.

Toliau, naudojant Gradient Boosting Machine modelį, pateikiame tuos pačius duomenų rinkinius ir atliekame analizę, naudojant tas pačias tris parametrų kombinacijas (kurias papildė mokymosi greičio parametras), siekiant nustatyti kuri iš kombinacijų leidžia pasiekti efektyviausius rodiklius tyrime. Žemiau 9 lentelėje pateikiame atlikto eksperimento rezultatus:

Naudojamo GBM modelio pirminių rezultatų analizė taikant skirtingus parametrus

Duomenų rinkinys (DR)	Modelis	Medžiai	Gylis	Stebėjimų skaičius	Mokymosi greitis	Klasių Balansas	Balanso dydis	AUC	PR_AUC	R ²	False Positives	FP rate	F1 Score
Pradinis duomenų rinkinys	GBM	15	5	10	0,01	1, 1	5	0,93	0,67	0,18	134	23,80%	0,71
Pradinis duomenų rinkinys	GBM	35	10	15	0,03	1, 3	5	0,94	0,68	0,33	106	19,00%	0,75
Pradinis duomenų rinkinys	GBM	50	7	20	0,05	1, 5	5	0,95	0,75	0,56	94	17,7%	0,75
Pradinis DR su procesų gavybos savybėmis	GBM	15	5	10	0,01	1, 1	5	0,92	0,69	0,33	92	16,7%	0,75
Pradinis DR su procesų gavybos savybėmis	GBM	35	10	15	0,03	1, 3	5	0,94	0,7	0,44	73	13,80%	0,76
Pradinis DR su procesų gavybos savybėmis	GBM	50	7	20	0,05	1, 5	5	0,95	0,75	0,57	59	13,9%	0,76
Pradinis DR su didelio kalbos modelio savybėmis	GBM	15	5	10	0,01	1, 1	5	0,93	0,58	0,16	134	25,20%	0,67
Pradinis DR su didelio kalbos modelio savybėmis	GBM	35	10	15	0,03	1, 3	5	0,89	0,63	0,3	119	22,30%	0,76
Pradinis DR su didelio kalbos modelio savybėmis	GBM	50	7	20	0,05	1, 5	5	0,92	0,67	0,4	65	14,4%	0,73
Pradinis DR su visomis savybėmis	GBM	15	5	10	0,01	1, 1	5	0,94	0,61	0,42	131	24,9%	0,67
Pradinis DR su visomis savybėmis	GBM	35	10	15	0,03	1, 3	5	0,95	0,71	0,56	67	14,00%	0,73
Pradinis DR su visomis savybėmis	GBM	50	7	20	0,05	1, 5	5	0,94	0,5	0,7	60	13,3%	0,74

Šaltinis: sudaryta autoriaus

Pagal 9 lentelėje pateiktus rezultatus matome, kad taikant pirminę parametrų kombinaciją, kurią susidaro 15 medžių, 5 mazgų gylio, 10 minimalių stebėjimų mazge, nustatant balanso dydžio reikšmę kaip 5 ir mokymosi greitį 0,01 gauname prasčiausius rezultatus, lyginant su kitomis kombinacijomis. Naudojant pirmąją parametrų kombinaciją modelis, kurią pritaikome skirtingiems duomenų rinkiniams, pastebime, kad modelis generuoja didžiausius klaidingai teigiamų įspėjimų kiekius, kurie varijuoja tarp 92 ir 134, o procentinė klaidingai teigiamų įspėjimų dalis nuo visų sugeneruojamų įspėjimų sudaro tarp 16,7% ir 24,9% (skirtinguose duomenų rinkiniuose). Taip pat ši parametrų kombinacija prasčiausiais kitais modelio bendriniais parametrais, pavyzdžiui, determinacijos koeficientas (R^2) varijuoja tarp 0,16 ir 0,42, F1 reikšmė varijuoja tarp 0,67 ir 0,75 PR_AUC kreivės ploto rodiklis grąžino reikšmes tarp 0,58 ir 0,69. Nepaisant to, kad naudojant pirmąją parametrų kombinaciją gauname mažiausias aukščiau išvardintų rodiklių reikšmes, tačiau pastebime, kad naudojant kitas parametrų kombinacijas šie rodikliai arba nežymiai padidėja arba šiek tiek sumažėja. Pagal tai darome išvadą, kad GBM modelis yra gana stabilus, sugeba panašiai gerai nustatyti retas klases, paaiškina tikslinio kintamojo variaciją, rodo pakankamai gerą tikslumo ir jautrumo balansą.

Tuo tarpu, taikant antrąją parametrų kombinaciją, kuri susidaro iš 35 medžių, 10 mazgų gylio, 15 minimalių stebėjimų mazge, padidinant klasių svorius iki 1:3 ir nustatant mokymosi greičio reikšmę 0,03 gauname pakankamai geresnius klaidingai teigiamų įspėjimų rezultatus visuose duomenų rinkiniuose. Šiame eksperimento etape naudojamas GBM modelis sugeneravo tarp 119 ir 67 klaidingai teigiamų įspėjimų skirtinguose duomenų rinkiniuose, o procentinė klaidingai teigiamų įspėjimų dalis nuo bendro įspėjimų kiekio varijuoja tarp 22,3% ir 14%. Kiti modelio rodikliai išliko panašūs kaip ir pirmosios parametrų kombinacijos metu. Galiausiai, taikant trečiąją parametrų kombinaciją, kurios metu didiname medžių kieki iki 50, minimalių stebėjimų mazge skaičių iki 20, klasių svorius iki 1:5 ir mokymosi greitį iki 0,05 o medžių gylį sumažinus iki 7, modelis pateikia dar geresnius klaidingai teigiamų įspėjimų rezultatus, kurie varijuoja tarp 94 ir 59 (skirtinguose duomenų rinkiniuose). Klaidingai teigiamų įspėjimų procentinė dalis nuo visų sugeneruotų įspėjimų dviejuose duomenų rinkiniuose sumažėja tarp 0,7% iki 1,3%, tačiau duomenų rinkinyje su procesų gavybos savybėmis padidėja 0,10%, o duomenų rinkinyje su didelio kalbos modelio savybėmis procentinė reikšmė sumažėja net 5,9%. Visi kiti modelio rodikliai išlieka panašūs į anksčiau gautus rezultatus.

Atlikę skirtingų modelių mašininio mokymosi procesą, kurio metu naudojome du skirtingus modelis, keturis skirtingus duomenų rinkinius, kurie buvo papildyti procesų gavybos ir didelio kalbos modelio savybėmis, atlikome eksperimentą, taikant skirtingus modelių parametrus. Šios eksperimento dalies metu nustatėme, kad nors Distributed Random Forest modelis geriausiai veikė pagal antrąją parametrų kombinaciją, tačiau naudojamas Gradient Boosting Machine pateikė dar geresnius klaidingai teigiamų įspėjimų rezultatus, naudojant trečiąją parametrų kombinaciją.

Įvertinus aukščiau pateiktus skirtingų taikomų parametrų rezultatus, išvedame galutinę parametrų kombinacijos versiją (kartu su kitais svarbiais parametrais). Žemiau pateikiamas konkretus parametrų sąrašas su papildomais paaiškinimais kaip parametrai daro įtaką modeliui:

1. **Response_column = isFraud** – sukčiavimo žymės kintamąjį pasirenkame kaip tikslinį (angl. label), kurį modelis turi prognozuoti, kadangi sukčiavimo aptikimo uždavinys yra klasifikacija.
2. **Ntrees = 50** – šiuo parametru nurodome, kiek sprendimų medžių naudojame mokymosi procese. Atlikto pirminio eksperimento metu ir kitais testavimo atvejais buvo pastebėta, kad didesni medžių kiekis nei 50 stipriai permoko modelį, kurį sudaro naudojamas algoritmas.
3. **Max_depth = 7** – šis parametras nustato maksimalų sprendimų medžio gylį. Pagal nustatytą reikšmę ribojame medžių gylį tam, kad modelis nesukurtų pernelyg sudėtingų sprendimų taisyklių, kurios yra paremtos nedideliais ir nestabiliais duomenų fragmentais.
4. **Min_rows = 20** – šis parametras nurodo minimalų stebėjimų skaičių, kuris yra reikalingas tolimesniam duomenų skaidymui. Naudodami nurodytą parametą užtikriname, kad kiekviena sprendimo dalis būtų paremta pakankamu duomenų kiekiu, o tai sumažina tikimybę, kad modelis išmoks atsitiktinius sukčiavimo atvejus.
5. **Learn_rate = 0.05** – mokymosi greičio parametras apibrėžia, kiek stipriai kiekvienas naujas modelio žingsnis koreguoja ankstesnes prognozes. Pasirinkta mažesnė nei standartinė reikšmė leidžia atlikti korekcijas palaipsniui, todėl padidiname mokymosi stabilumą ir sumažiname riziką, kad modelis per greitai prisitaikys prie specifinių ir nereikšmingų atvejų.
6. **Balance_classes = true** – kadangi sukčiavimo atvejai sudaro tik nedidelę visų mokėjimų dalį (0.3 %), todėl be papildomų priemonių modelis būtų linkęs prognozuoti tik dažniau pasitaikančią klasę. Pritaikant klasių balanso parametą, mokymo metu didesnis dėmesys suteikiamas retesnei klasei (sukčiavimo), todėl modelis mokosi atpažinti tiek įprastinius mokėjimus, tiek su sukčiavimu susijusias finansines operacijas.
7. **Class_sampling_factors = 1, 5** – šiuo parametru tiksliau valdome klasių balansavimo stiprumą, suteikdami papildomus svorius kategorinėms reikšmėms. Šiuo atveju įprastiniams mokėjimams taikoma 1 svorio reikšmė, o tuo tarpu sukčiavimo atvejams 5 kartus didesnė svorio reikšmė. Taikant šį parametą nurodome modeliui sukčiavimo kategorijos svarbą, todėl modelis jautriau reaguoja į šias reikšmes ir sumažina praleistų sukčiavimo atvejų skaičių.
8. **Max_after_balance_size = 5** – taikydami balanso dydžio parametą valdome duomenų padidėjimą po klasių balansavimo. Šiuo konkrečiu parametru nustatome, kad subalansuotas duomenų rinkinys negali būti daugiau nei penkis kartus didesnis už pradinį. Tokiu būdu išvengiame pernelyg didelio dirbtinio duomenų išplėtimo, kas gali pabloginti modelio stabilumą.

Pagal aukščiau pateiktus parametrus (parametrų naudojimo H2O AutoML pavyzdžiai pateikiami 6 ir 7 prieduose) atliekame pakartotinį mašininio mokymosi procesą. Šio proceso metu naudojami Distributed Random Forest ir Gradient Boosting Machine modeliai ir jiems pateikiami aukščiau minėti finansinių duomenų rinkiniai, papildyti įvairiomis gautomis savybėmis iš procesų gavybos ir didelio kalbos modelio. Taip pat atliekant mašininio mokymosi procesą ir siekiant išvengti nereikalingo triukšmo, kuris galėtų iškraipyti ar pabloginti galutinius rezultatus, modeliai formuojami išėmus šiuos kintamuosius: transDateAndTime, customerNumber, merchantCleaned, firstName, lastName, street, city, zipCode, customerLatitude, customerLongitude, transNum, unixTim, merchantLatitude, merchantLongitude (pavyzdys pateikiamas 8 priede). Šis sprendimas yra priimamas atsižvelgiant į mažą kintamųjų reikšmingumą likusių kintamųjų atžvilgiu. Verta paminėti, kad šie kintamieji buvo išimti ir anksčiau atlikto eksperimento metu, kuomet buvo taikomos skirtingos parametrų kombinacijos. Žemiau 10 lentelėje pateikiami abiejų modelių galutiniai gauti rezultatai iš mokymosi proceso:

Taikomo metodo metu išgautų savybių galutiniai efektyvumo rezultatai

Duomenų rinkinys (DR)	Modelis	Medžiai	Gylis	Stebėjimų skaičius	Mokymosi greitis	Klasių Balansas	Balanso dydis	AUC	PR_AUC	R ²	False Positives	FP rate	F1 Score
Pradinis duomenų rinkinys	DRF	50	7	20	N/A	1, 5	5	0,97	0,63	0,38	125	24,30%	0,67
Pradinis duomenų rinkinys	GBM	50	7	20	0,05	1, 5	5	0,95	0,75	0,56	94	17,7%	0,75
Pradinis DR su procesų gavybos savybėmis	DRF	50	7	20	N/A	1, 5	5	0,97	0,69	0,3	114	20,80%	0,71
Pradinis DR su procesų gavybos savybėmis	GBM	50	7	20	0,05	1, 5	5	0,95	0,75	0,57	59	13,9%	0,76
Pradinis DR su didelio kalbos modelio savybėmis	DRF	50	7	20	N/A	1, 5	5	0,98	0,69	0,43	118	23%	0,67
Pradinis DR su didelio kalbos modelio savybėmis	GBM	50	7	20	0,05	1, 5	5	0,92	0,67	0,4	65	14,4%	0,73
Pradinis DR su visomis savybėmis	DRF	50	7	20	N/A	1, 5	5	0,98	0,73	0,43	98	18,50%	0,72
Pradinis DR su visomis savybėmis	GBM	50	7	20	0,05	1, 5	5	0,94	0,5	0,7	60	13,3%	0,74

Šaltinis: sudaryta autoriaus

Pagal 10 lentelėje pateiktus rezultatus, matome, kad taikant tuos pačius modelio parametrus, Gradient Boosting Machine modelis pateikia geresnius klaidingai teigiamų įspėjimų rezultatus. Pavyzdžiui, Distributed Random Forest gaunami klaidingai teigiamų įspėjimų skaičius per keturis skirtingus duomenų rinkinius varijuoja tarp 98 ir 125, o tuo tarpu Gradient Boosting Machine modelio generuojamų klaidingai teigiamų įspėjimų skaičius varijuoja tarp 59 ir 94. Taip pat DRF modelio kiti parametrai, determinacijos koeficientas (R^2), PR_AUC kreivės plotas, F1 rodiklis, yra žymiai prastesni visuose duomenų rinkiniuose, apart duomenų rinkinio, kuris yra sudarytas iš pradinų duomenų ir didelio kalbos modelio savybių (rodikliai nežymiai geresni). Tarp skirtingų modelių ir skirtingų duomenų rinkinių matome panašią klaidingai teigiamų įspėjimų generavimo tendenciją, kuomet naudojant pradinį duomenų rinkinį tokių įspėjimų skaičius didžiausias, o reikšmingiausias sumažėjimas matomas naudojant pradinį duomenų rinkinį su procesų gavybos savybėmis ir naudojant pradinį duomenų rinkinį su procesų gavybos ir didelio kalbos modelio savybėmis.

Kadangi Gradient Boosting Machine gauti rezultatai yra reikšmingai geresni, o bendra vertinimo tarp modelių tendencija yra itin panaši, todėl toliau atliekame detalesnę gautų rezultatų tik iš GBM modelio analizę, kurios metu galime išsamiai įvertinti, kaip skirtingi savybių rinkiniai daro įtaką modelio gebėjimui sumažinti klaidingai teigiamų sprendimų skaičių. Naudojant tik pradinį duomenų rinkinį, modelis pasiekia aukštą atskyrimo gebą ($AUC = 0.95$), tačiau sugeneruoja didžiausią klaidingai teigiamų įspėjimų kiekį – 94, o klaidingai teigiamų perspėjimų santykis su visais perspėjimais sudaro net 17,7 proc. Į duomenų rinkinį įtraukus savybes, išgautas taikant proceso gavybos metodus, pastebime, kad klaidingai teigiamų įspėjimų kiekis sumažėjo iki 59 (procentine dalimi – 37,2 proc. sumažėjimas), o klaidingai teigiamų įspėjimų santykis su visų įspėjimų kiekiu sumažėjo iki 13,9 proc. Taip pat pagal kitus turimus rodiklius matome (AUC ir PR_AUC), kad modelis pateikia geresnius rezultatus, nors bendras modelio tikslumas išliko toks pats.

Tuo tarpu, naudojant tik didelio kalbos modelio (gpt-4.1-mini) sugeneruotus įvėčius (kuriuos laikome kaip papildomas savybes), pastebime šiek tiek kitokią tendenciją. Nors šį duomenų rinkinį naudojant modelis sugeneravo mažesnę klaidingai teigiamų įspėjimų kiekį (65), o klaidingai teigiamų įspėjimų kiekio santykis su visais sugeneruotais perspėjimais sudaro 14,4 proc., tačiau šie rodikliai yra prastesni nei modelyje su proceso gavybos savybėmis. Tu pačiu matome, kad šio modelio bendrosios atskyrimo gebos metrikos ($AUC = 0.92$, PR_AUC = 0.67) yra prastesnės už anksčiau įvertintus modelius, o tai reiškia, kad naudojant tik didelio kalbos modelio įvėčius kaip papildomas savybes, sukuria modelį, kuris nepakankamai stabiliai atspindi sukčiavimo struktūrinius dėsningumus. Galiausiai, sudarytas modelis, naudojant pradinį duomenų rinkinį, kuris yra papildytas proceso gavybos ir didelio kalbos modelio naudojimo etapuose gautomis savybėmis, parodė panašius rezultatus kaip ir savybės iš procesų gavybos etapo. Modelis pagal naudojamą duomenų rinkinį

sugeneravo 60 klaidingai teigiamų įspėjimų, o klaidingai teigiamų įspėjimų santykis su visais sugeneruotais įspėjimais siekia 13,3 proc. Nors kitos bendrinės metrikos yra mažesnės už pirmus du modelius, tačiau praktiniu požiūriu ketvirtasis modelis gali būti laikomas panašiai efektyviu kaip procesų gavybos savybėmis paremtas antrasis modelis. Taip pat verta paminėti, kad klaidingai teigiamų įspėjimų santykinis kiekis su visais sugeneruotais įspėjimais yra žemiausias iš visų pateiktų modelio variantų (13,3 proc.). Galiausiai, galime teigti, kad didelio kalbos modelio įverčiai nedaro reikšmingos įtakos, siekiant pagerinti modelio efektyvumą, mažinant klaidingai teigiamų įspėjimų kiekį.

Galima atlikti eksperimentinę darbo dalį ir gavus galutinius siūlomo metodo rezultatus, atliktame metodo efektyvumo palyginimą su kitais siūlomais metodais, kurių analizė buvo atliekama literatūros analizės dalyje. Žemiau pateikiama 11 lentelė, kurioje siūlomo metodo rezultatus lyginame su kitais penkiais metodais, naudojamais mokėjimų stebėsenos ir sukčiavimo prevencijos kontekste:

11 lentelė

Taikomo metodo efektyvumo palyginimas su kitais mokslinėje literatūroje siūlomais metodais

Metodas	AUC	PR_AUC	R ²	False Positives	FP rate	F1 Score
Procesų gavyba ir dideli kalbos modeliai	0,94	0,5	0,7	60	13,3%	0,74
GAN-VAE anomalijų aptikimo modelis	N/A	N/A	N/A	N/A	N/A	0,795
GBM + dideli kalbos modeliai	0,89	N/A	N/A	N/A	N/A	N/A
Didelis kalbos modelis (In Context Learning)	N/A	N/A	N/A	N/A	N/A	0,651
Dideli kalbos modeliai AML užduotims	N/A	N/A	N/A	0	0%	1
Laikinis grafų neuroninis tinklas (TGNN)	0,987	N/A	N/A	N/A	-37%	0,93

Šaltinis: sudaryta autoriaus pagal atlikto eksperimento rezultatus ir literatūros analizės dalyje atliktų mokslinių tyrimų empirinius rezultatus

Pagal pateiktą susistemintą siūlomo metodo ir kitų esamų metodų rezultatų lentelę matome, kad pateikiamų rezultatų duomenys yra nevientisi, o tai pažymi esamą problemą atliktų mokslinių darbų kontekste. Iš vertinamų metodų, tik didelių kalbos modelių taikymo AML užduotims autorius pateikia realius klaidingai teigiamų įspėjimų skaičius, tačiau dėl itin mažos imties (10 mokėjimų) tokio rezultato negalime laikyti reprezentatyviu. Didžioji dalis metodų pateikia F1 rodiklį, kuris nurodo kaip tiksliai siūlomų metodų modeliai gali aptikti teisingus sukčiavimo atvejus, išvengiant perteklinio klaidingai teigiamų įspėjimų kiekio. Pagal šį rodiklį, darbe siūlomą procesų gavybos ir didelių kalbos modelių metodą lenkia du siūlomi metodai – GAN-VAE anomalijų aptikimo modelis ir Laikinis grafų neuroninis tinklas (TGNN). Galiausiai, būtina pabrėžti, kad visi aukščiau išvardinti metodai naudoja skirtingus finansinių duomenų rinkinius, mokėjimų požymius ar net skirtingai apibrėžia siūlomo

metodo užduotį, todėl gaunami rezultatai nėra vientisi, todėl tarpusavio lyginimas tampa nebeobjektyvus.

Taip pat verta paminėti, kad siūlomo metodo efektyvumo analizė yra vykdoma pagal taikomą savybių inžinerijos procesą, t.y. naudojant pirminį finansinių duomenų rinkinį, taikant procesų gavybą ir didelį kalbos modelį, išvedame papildomas savybes, kurios gali padėti efektyviau nustatyti sukčiavimo atvejus, taip mažinant klaidingai teigiamų įspėjimų kiekį. Eksperimento metu atliekant savybių inžineriją, prie pradinio duomenų rinkinio pridedamos skirtingos savybių grupės (tik procesų gavybos savybės, tik didelio kalbos modelio savybės ir procesų bei didelio kalbos modelio savybės kartu), todėl tolimesnės analizės, naudojant H2O AutoML įrankį, metu galime analizuoti ar papildomos išvestos savybės gali padėti sumažinti modelio generuojamą klaidingai teigiamų įspėjimų kiekį.

IŠVADOS

1. Pirmoje darbo dalyje atlikta pinigų plovimo sampratos analizė ir pagrindinių veikimo etapų – perkėlimo, sluoksniavimo ir integracijos – suvokimas, leidžia geriau suprasti, šio proceso prevencinę svarbą finansų įstaigose. Analizuodami pinigų plovimo etapus suprantame šio proceso prevencijos svarbą, kuri yra išreiškiama tiek nacionalinių priežiūros institucijų (pavyzdžiui, Lietuvos Bankas) lūkesčiais ir reikalavimais, tiek nacionalinio ir tarptautinio lygmens teisės aktais, įstatymais ir direktyvomis. Dėl šių priežasčių suprantame, kad finansų įstaigos yra orientuotos į ne tik efektyvų klientų aptarnavimą, tačiau skiria didelį dėmesį ir išteklius kliento pažinimo procesams, sandorių stebėsenai, rizikos vertinimui, informacijos saugojimui ir vidaus kontrolės mechanizmams. Nagrinėjami tiek nacionalinio, tiek tarptautinio reguliavimo aspektai leidžia geriau suprasti, kokio sisteminio požiūrio reikalauja pinigų plovimo prevencijai ir kokie praktiniai veiksmai taikomi šioms rizikoms suvaldyti finansų sektoriuje. Technologiškai lanksti infrastruktūra yra svarbi, siekiant realiuoju laiku blokuoti aukštos rizikos operacijos, atlikti išsamią retrospektyvią mokėjimų analizę ir pritaikyti segmentinius stebėjimo scenarijus pagal rizikos požymius. Tuo pačiu, sistema turi būti pajėgi stebėti neaktyvias ar tranzitines sąskaitas, atsižvelgti į kitų procedūrų metu surenkamą informaciją. Atsižvelgiant į šiuos reikalavimus, galime teigti, kad akademiniėje literatūroje vyrauja panašus požiūris, kad taisyklėmis grįsta stebėsenos sistema nebeatitinka keliamų lūkesčių ir rinkos standartų, o generuojamų klaidingų įspėjimų kiekis ženkliai didina institucijos žmogiškųjų ir finansinių išteklių kaštus. Dėl priklausomybės nuo žmogiškojo faktoriaus ir negebėjimo prisitaikyti prie naujų pinigų plovimo schemų, taisyklėmis grįstą sistemą privalo pakeisti pažangesni metodai, kurie stiprintų stebėsenos infrastruktūrą, taupytų įstaigos išteklius ir gerintų efektyvumo rodiklius. Dėl šių priežasčių darome išvadą, kad mokėjimų stebėsenos taikymo metodai turi būti nuolat atnaujinami, tiriant esamas ir inovatyvias technologijas, kurios leistų pagerinti atsparumą sukčiavimui, pinigų plovimui, mažintų klaidingai įspėjimų kiekį.
2. Siekiant atliepti technologinio inovatyvumo trūkumą dabartinėse, taisyklėmis grįstose, mokėjimų stebėsenos sistemose, akademiniai šaltiniai nagrinėja įvairių algoritmų ir kitų technologijų panaudojimo būdus, siekiant tiksliai ir efektyviai identifikuoti sukčiavimo atvejus finansinių mokėjimų duomenų masyvuose. Viena iš siūlomų sprendimų, galinčių pasiūlyti efektyvesnį sprendimą yra procesų gavybos algoritmai, kurie geba iš įvykių žurnalų atkurti įvykusių procesų modelius, nustatyti nukrypimus ir identifikuoti atsiradusias ribines reikšmes ir anomalijas, kurios gali padėti identifikuoti nusikalstamą finansinę veiklą. Atlikta mokslinių šaltinių, kurie atlieka empirinius tyrimus, darome išvadą, kad dažniausiai ir

efektyviausiai yra taikomas neapibrėžtas algoritmas (angl. Fuzzy miner) dėl savo gebėjimo filtruoti triukšmą ir pateikti aiškiai suprantamą modelį, vengiant „spaghetti-like“ problematikos. Euristinis (angl. Heuristic miner) pasižymi savo tikslumu ir tinkamumu dirbti su triukšmingais įvykių žurnalais, todėl yra tinkamas kliento pažinimo ir auditų procesuose. Tuo tarpu indukcinis algoritmas (angl. Inductive miner) taikomas mažiausiai dėl savo jautrumo triukšmui ir sudėtingų modelių generavimo, todėl nėra tinkamas efektyviai analizei. Atliekami eksperimentai siūlo taikyti principą, kuomet sukčiavimo prevencijos priemonė/metodas yra taikomas naudojant procesų gavybos elementą ir kitą duomenų analizės metodą/techniką, kadangi procesų gavybos algoritmai turi tam tikrą technologišią ribotumą ir reikalauja didesnio žmonių įsikišimo.

Priėjus išvadą, kad procesų gavybos algoritmai veikia efektyviau, naudojant juos su kitomis duomenų analizės metodais ir technologijomis, toliau atliekame literatūros, nagrinėjančios didelių kalbos modelių pritaikomumo mokėjimų stebėsenoje, šaltinius. Šios inovatyvios technologijos panaudojimas mokėjimų stebėsenos metoduose/sistemos atliepia poreikį skatinti šios procedūros modernumą ir efektyvumą, kurio būtinybę akcentuoja moksliniai darbai. Remiantis akademiniuose darbuose atliktų tyrimų empiriniais duomenimis, didelių kalbų modeliai, tokie kaip GPT, BERT, RoBERTa ar BloombergGPT demonstruoja pakankamai aukštą prognozavimo efektyvumą (F1 score rodiklis varijuoja nuo 0,65 iki 1,00), geba analizuoti tiek struktūrizuotus, tiek nestruktūrizuotus duomenis, įvertinti semantinį kontekstą bei generuoti paaiškinimus ir rekomendacijas, padedančias analitikams priimti pagrįstus sprendimus. Panašius rezultatus generuoja ir kiti dirbtinio intelekto modeliai (grafiniai neuro-tinklai, Gradient Boosting Machine ir kt.), kurie neretai būna derinami kartu su didelių kalbų modeliais. Vis dėlto, šių technologijų praktinis taikymas susiduria su įvairiais iššūkiais – didelės skaičiavimo sąnaudos, duomenų saugumas/privatumas, algoritmo šališkumas, galimos modelių haliucinacijos ir kt. – todėl inovatyvios technologijos taikymas privalo būti prižiūrimas kompetentingų specialistų, kurie galėtų užtikrinti rezultatų patikimumą ir atitiktį reguliaciniams reikalavimams. Nepaisant to, apibendrinus galime daryti išvadą, kad didelių kalbos modelių technologija turi pakankamą potencialą transformuoti mokėjimų stebėsenos sistemų veikimo principą, didinant jų tikslumą, efektyvumą bei geresnį pritaikomumą nuolat kintančiai aplinkai, t.y. sudėtingėjant mokėjimų srautų struktūroms, griežtėjant atitikties reikalavimams ir pan.

3. Remiantis atlikta literatūros analizės dalimi, antroje darbo dalyje yra siūlomas hibridinis procesų gavybos neapibrėžto algoritmo (angl. Fuzzy Miner) ir didelio kalbos modelio (šiuo atveju gpt-4.1-mini) metodas ir jo taikymo būdas, siekiant atlikti efektyvią pirminių finansinių mokėjimų duomenų rinkinio analizę ir pagal tai išvesti papildomas savybes (angl. features),

kurios gali padėti sumažinti klaidingai teigiamų išpėjimų kiekį. Toks taikomo metodo principas leidžia efektyviau nustatyti naujausias sukčiavimo ir pinigų plovimo tendencijas, kadangi generuojamos savybės gali būti kintančios, priklausomai nuo identifikuotų įtartinų finansinės elgsenos požymių. Taip pat didelio kalbos modelio technologinis pranašumas leidžia įvertinti ne tik skaitinę, bet ir tekstinę informaciją, todėl didesnę prasmę įgauna ir kitų procesų metu renkama tekstinė informacija – mokėjimo paskirties laukai, kliento pažinimo profiliai ir t.t. Procesų gavybos etapui pasirinkimas taikyti Fuzzy Miner algoritmą yra pagrįstas finansinių duomenų kompleksiskumu ir apimtimi, todėl teigiame, kad būtent šis algoritmas gali veikti efektyviausiai nagrinėjant tokios struktūros duomenų masyvus. Pasitelkiant Celonis aplinkos siūlomus analitinius įrankius – Process Explorer ir Case Explorer – atliekame išsamią proceso ir atskirų atvejų analizes. Šis funkcionalumas yra svarbus, siekiant rasti priežastingumo ryšius, kurie dažniausiai lemia sukčiavimo ir pinigų plovimo įvykius. Identifikuoti priežastingumo ryšiai yra išreiškiami naujai sukurtais savybėmis. Tuo tarpu paraleliai naudojamas GPT modelis (gpt-4.1-mini) papildo procesų gavybos etapą semantinio vertinimo aspektu. Taikant mokymosi be pavyzdžių (angl. Zero-shot) metodą ir papildomai sugeneruotus kliento rizikos profilius, atliekame visų mokėjimo operacijų semantinį ir bendrą rizikos vertinimą, kurio metu gaunami modelio įverčiai naudojami kaip papildomos savybės, galinčios sumažinti klaidingai teigiamų išpėjimų kiekį. Papildomai sąveikaujant su modeliu per API sąsaja atsiranda galimybė koreguoti modelio parametrus, jų jautrumą, kontroliuoti gaunamus rezultatus, jų pateikimo formą ir kitus svarbius aspektus, kurie yra svarbūs mokėjimų stebėsenos procese. Ši siūlomo metodo dalis suteikia papildomą interpretacinę dimensiją, kurios klasikiniai statistiniai ir taisyklėmis paremti metodai negali suteikti.

4. Galiausiai, remiantis pasiūlyto metodo veikimo principu ir nurodyta veikimo schema atliekame pirminį duomenų paruošimą, o paruoštą duomenų rinkinį su papildomais išvestais kintamaisiais paraleliai įkeliamo tiek į procesų gavybos platformą, tiek per API pateikiame naudojamam dideliame kalbos modeliui. Atlikus šį etapą iš procesų gavybos analizės ir didelio kalbos modelio vertinimų išvedame naujas savybes, kurios buvo identifikuotos ar įvertintos analizės metu. Galiausiai sudarome bendrai keturis pirminio duomenų rinkinio variantus, t.y. tik pirminiai duomenys, pirminiai duomenys ir savybės iš procesų gavybos, pirminiai duomenys ir įverčiai (kuriuos laikome kaip savybėmis) iš didelio kalbos modelio bei pirminiai duomenys tiek su procesų gavybos, tiek su didelio kalbos modelio savybėmis.

Visus šiuos keturis duomenų rinkinius įkeliamo į H2O AutoML automatizuotą mašininio mokymosi platformą ir pagal darbe pateiktas parametrų kombinacijas naudojame Distributed Random Forest (DRF) ir Gradient Boosting Machine (GBM) modelius, siekiant

nustatyti optimaliausią parametrų kombinaciją, pagal kurią modelis generuotų geriausius rezultatus (mažintų klaidingai teigiamų įspėjimų skaičių). Atlikto eksperimento metu išvedame pagrindinę parametrų kombinaciją ir naudojant išskirtus duomenų rinkinius atliekame vertinimą, kaip pateikiami duomenų rinkiniai (su skirtingomis savybėmis) daro įtaką modelių galutiniams rezultatams (klaidingai teigiamų įspėjimų mažinimas). Atliekamo eksperimento eigoje nustatėme, kad naudojantis ta pačia parametrų kombinacija Gradient Boosting Machine (GBM) modelis pateikė reikšmingai geresnius rezultatus beveik visuose duomenų rinkiniuose, o kadangi abiejų modelių rezultatai skirtinguose duomenų rinkiniuose parodė vienodą tendenciją, todėl galutiniam etape pasirenkame pagrinde vertinti GBM pateiktus rezultatus.

Susisteminius gautų rezultatų duomenis, galime daryti išvadą, kad santykinai didžiausią teigiamą poveikį klaidingai teigiamų atvejų mažinimui turi sudarytos naujos savybės iš procesų gavybos etapo, kuris lyginant su pirminiu duomenų rinkiniu sumažino tokių įspėjimų kiekį 37,2 proc. Tuo tarpu iš didelio kalbos modelio (gpt-4.1-mini) gauti įverčiai, kuriuos laikome papildomomis savybėmis parodė šiek tiek prastesnę rezultatą, lyginant su procesų gavybos savybėmis pradinio duomenų rinkinio atžvilgiu, t.y. klaidingai teigiamų įspėjimų kiekis sumažėja 30,8 proc. Tuo tarpu, vertindami modelio rezultatus pagal pateiktą duomenų rinkinį, sudarytą tiek iš procesų gavybos, tiek iš didelio kalbos modelio savybių, matome, kad klaidingai teigiamų įspėjimų kiekis (lyginant su pirmuoju duomenų rinkiniu) sumažėja 36,1 proc., t.y. 1,1 proc. prasčiau nei taikant tik proceso gavybos savybes.. Pagal gautus modelio rezultatus naudojant ketvirtąjį duomenų rinkinį, galime teigti, kad neigiamą įtaką modelio rezultatams galėjo padaryti didelio kalbos modelio savybės, kadangi jų efektyvumas yra reikšmingai mažesnis nei proceso gavybos savybių. Taip pat pastebime, kad ketvirtojo modelio bendrinės metrikos (AUC, PR_AUC, F1 Score) yra prastesnės nei naudojant kitus duomenų rinkinius, tačiau klaidingai teigiamų įspėjimų santykinė dalis su visais sugeneruotais įspėjimais yra mažiausia (13,3 proc.). Apibendrinus gautus eksperimento rezultatus drįstame teigti, kad naudojant siūlomą procesų gavybos ir didelio kalbos modelio metodą, galime reikšmingai sumažinti klaidingai teigiamų įspėjimų kiekį daugiau nei trečdaliu (36,1 proc.), todėl toks rezultatas galėtų reikšmingai padidinti sistemos efektyvumą ir sumažinti operacinius kaštus bei efektyviau panaudoti turimus žmogiškuosius išteklius.

PASIŪLYMAI

1. Siūloma procesų gavybos etapą atlikti naudojant kitas platformas/įrankius, kurios yra ne komercinio, o akademinio tipo ir yra plačiai naudojamos moksliniuose tyrimuose, pavyzdžiui ProM Framework, PM4Py (Process Mining for Python) ar Aprome. Šie įrankiai yra labiau universalūs ir turi daugiau procesų gavybos algoritmų, kuriuos galima taikyti atliekant siūlomo metodo eksperimentus. Taip pat akademiniai įrankiai turi daugiau reguliuojamų parametrų, kurie gali padėti dar labiau padidinti analizės efektyvumą, siekiant išgauti kuo geresnės kokybės savybes.
2. Siūlomo eksperimento metu, naudodami didelį kalbos modelį taikome mokymosi be pavyzdžių (angl. Zero-Shot) metodą, kuris laikomas vienu iš paprasčiausių, atliekant užduotis. Kadangi sukčiavimo ir pinigų plovimo elgsena turi tam tikrus pasikartojančius modelius/etapus, todėl verta pabandyti taikyti kitus metodus, kurie iš anksto paruošia didelį kalbos modelį užduoties atlikimui, pavyzdžiui mokymąsi su minimaliu pavyzdžių kiekiu (angl. Few-Shot), išankstinį apmokymą (angl. Pre-training), instrukcijomis grįstą apmokymą (angl. Instruction Tuning). Šiuos mokymosi metodus verta išbandyti, kadangi empirinių rezultatų analizė parodė, kad didelio kalbos modelio įverčiai, kurie yra naudojami kaip papildomos savybės yra reikšmingai mažiau efektyvios nei procesų gavybos metu išgaunamos savybės, todėl kombinuojant šių dviejų procesų rezultatus, galutiniai rezultatai tampa prastesni nei taikant juos atskirai.
3. Eksperimento metu buvo naudojami du modeliai H2O AutoML platformoje, t.y. Distributed Random Forest ir Gradient Boosting Machine, todėl tolimesniuose tyrimo etapuose siūloma į eksperimento dalį įtraukti daugiau modelių, kurie veikia kaip klasifikatoriai, pavyzdžiui, Generalized Linear Modeling (GLM), Naive Bayes, RuleFit ir kt.
4. Taip pat siūloma atlikti papildomus empirinius tyrimus, taikant siūlomą procesų gavybos ir didelio kalbos modelio metodą, naudojant kitokio tipo ir apimties finansinių mokėjimų duomenis. Dabartiniai mokėjimai apima tik vienos valstybės ir tik debetinių kortelių mokėjimus, tačiau praktikoje įvairios finansų įstaigos siūlo didelį pasirinkimą įvairių mokėjimo produktų, kuriuos gali naudoti įvairaus tipo klientai (institucijos, juridiniai asmenys, fiziniai asmenys). Didesnės apimties ir įvairesni duomenys iš mokėjimo metodų perspektyvos (pvz., SWIFT tarptautiniai mokėjimai, SEPA ir pan.) galėtų parodyti tikslesni metodo efektyvumą, ypač vykdant mokėjimus tarp juridinių asmenų, tarp kurių pinigų plovimas vyksta dažniau ir įvairesnėmis strategijomis.

LITERATŪROS SĄRAŠAS

1. Aalst, W. M. (2022). Foundations of Process Discovery. Esantis W. M. Aalst, *Process Mining Handbook* (p. 37-76). Springer Nature.
2. Aalst, W. M. (2022). Process Mining: A 360 Degree Overview. Esantis W. M. Aalst, & J. Carmona, *Process Mining Handbook* (p. 3-37). Springer Nature.
3. Accorsi, R., & Leberherz, J. (2022). A Practitioner's View on Process Mining Adoption, Event Log Engineering and Data Challenged. Esantis W. M. Aalst, & J. Carmona, *Process Mining Handbook* (p. 212-243). Springer Nature.
4. Baader, G., & Krcmar, H. (2018). Reducing false positives in fraud detection: Combining the red flag approach with process mining. *International Journal of Accounting Information Systems*, 31, 1-16.
5. Carmona, J. (2010). Projection approaches to process mining using region-based techniques. *Data Mining and Knowledge Discovery*, 24(1), 218-246.
6. Carmona, J., Dongen, B. v., & Weidlich, M. (2022). Conformance Checking: Foundations, Milestones and Challenges. Esantis W. M. Aalst, & J. Carmona, *Process Mining Handbook* (p. 155-193). Springer Nature.
7. Celik, U., & Akcetin, E. (2018). Process mining tools comparison. *Online Academic Journal of Information Technology*, 9(34), 97-104.
8. Detten, J. N., Schumacher, P., & Leemans, S. J. (2023). An Approximate Inductive Miner. *5th International Conference on Process Mining (ICPM)* (p. 129-136). IEEE.
9. European Parliament & Council. (2018). Directive (EU) 2018/1673 of the European Parliament and of the Council of 23 October 2018 on combating money laundering by criminal law. Nuskaityta iš <https://eur-lex.europa.eu/eli/dir/2018/1673/oj/eng>
10. Fan, M. (2024). LLMs in Banking: Applications, Challenges, and Approaches. *In Proceedings of the International Conference on Digital Economy, Blockchain and Artificial Intelligence*, (p. 314-321).
11. Ferraris, A. F., Audrito, D., Caro, L. D., & Poncibo, C. (2025). The architecture of language: Understanding the mechanics behind LLMs. *In Cambridge Forum on AI: Law and Governance*, 1, 1-19.
12. Financial Action Task Force. (2018). *Professional Money Laundering*. FATF.
13. Garcia, S., Luengo, J., & Herrera, F. (2015). Data preprocessing in data mining. *Intelligent Systems Reference Library*, 72, 59-139.

14. Glimour, P. O. (2024). *Reexamining the 'Placement-Layering-Integration' Model of Money Laundering*.
15. Gunther, C. W., & Aalst, W. M. (2007). Fuzzy mining—adaptive process simplification based on multi-perspective metrics. *International conference on business process management* (p. 328-343). Springer Berlin Heidelberg.
16. Gupta, E. P. (2014). Process Mining A Comparative Study. *International Journal of Advance Research in Computer and Communication Engineering*, 3(11), 8594-8598.
17. Gurram, A. (2025). Transforming Regulatory Compliance with Generative AI. *Journal of Computer Science and Technology Studies*, 7(8), 1089-1096.
18. Han, J., Huang, Y., & Liu, S. T. (2020). Artificial Intelligence for Anti-Money Launderng - A Review and Extension. *Digital Finance*, 2(3), 211-239.
19. Yan, Y., Hu, T., & Zhu, W. (2024). Leveraging large language models for enhancing financial compliance: A focus on anti-money laundering applications. *4th International Conference on Robotics, Automation and Artificial Intelligence (RAAI)* (p. 260-273). IEEE.
20. Ye, M. (2007). A framework for data mining-based anti-money laundering research. *Journal of Money Laundering Control*, 10(2), 170-179.
21. Yin, F., Cao, J., Chu, J., & Olga, S. (2021). Warehousing Process Mining Research Based on Petri Net. *The International Conference on Artificial Intelligence and Logistics Engineering* (p. 54-65). Springer International Publishing. doi:https://doi.org/10.1007/978-3-030-80475-6_6
22. Jans, M., Werf, J. M., Lybaert, N., & Vanhoof, K. (2011). A business process mining application for internal transaction fraud mitigation. *Expert Systems with Applications*, 38(10), 13351-13359.
23. Ju, C., Ma, X., & Dong, B. (2025). Real-time Cross-border Payment Fraud Detection Using Temporal Graph Neural Networks: A Deep Learning Approach. *Academic Journal of Sociology and Management*, 1-12.
24. Ketenci, U. G., Kurt, T., Onal, S., Erbil, C., Akturkoglu, S., & Ilhan, H. S. (2021). A time-frequency based suspicious activity detection for anti-money laundering. *IEEE Access*, 9, 59957-59967.
25. Korkanti, S. (2024). Enhancing Financial Fraud Detection Using LLMs and Advanced Data Analytics. *2nd International Conference on Self Sustainable Artificial Intelligence Systems (ICSSAS)* (p. 1328-1334). IEEE.
26. Leoni, M. d. (2022). Foundations of Process Enhancement. Esantis W. M. Aalst, & J. Carmona, *Process Mining Handbook* (p. 243-274). Springer Nature.
27. Li, Y. (2023). A practical survey on zero-shot prompt design for in-context learning.

28. Lietuvos Bankas. (2022). *Nuolatinės kliento dalykinių santykių ir operacijų (sandorių) stebėsenos apžvalga*. Finansinių paslaugų ir rinkų priežiūros departamentas. Lietuvos Bankas.
29. Marin-Castro, H. M., & Tello-Leal, E. (2021). Event Log Preprocessing for Process Mining: A Review. *Applied Sciences*, 11(22), 10556.
30. Marvin, G., Hellen, N., Jjingo, D., & Nakatumba-Nabende, J. (2023). Prompt Engineering in Large Language Models. *International conference on data intelligence and cognitive informatics* (p. 387-402). Springer Nature.
31. Mittapelly, A. K., Tarra, V. K., & Agrahar, S. R. (2024). Autonomous CRM Agents: Using LLMS for Loan Processing, KYC and Regulatory Adherence in Salesforce. *Journal of Artificial Intelligence, Machine Learning and Data Science*, 2(3), 2624-2631.
32. Mulig, E. V., & Smith, M. (2004). Understanding and preventing money laundering. *Internal auditing*, 19(5), 22-25.
33. Nai, R., Sulis, E., Audrito, D., Trifiletti, V. M., Meo, R., & Genga, L. (2025). Leveraging process mining and event log enrichment in European public procurement analysis: a case study. *Computer Law & Security Review*, 57, 1-12.
34. Naik, P. V., Dintakurthi, N. K., Hu, Z., Wang, Y., & Qiu, R. (2025). Co-Investigator AI: The Rise of Agentic AI for Smarter, Trustworthy AML Compliance Narratives.
35. Naveed, H., Khan, A. U., Qiu, S., Saqib, M., Anwar, S., Usman, M., . . . Mian, A. (2025). A comprehensive overview of large language models. *ACM Transactions on Intelligent Systems and Technology*, 16(5), 1-72.
36. Olateju, O. O., Okon, S. U., Igwenagu, U. T., Salami, A. A., Oladoyinbo, T. O., & Olaniyi, O. O. (2024). Combating the Challenges of False Positives in AI-Driven Anomaly Detection Systems and Enhancing Data Security in the Cloud. *Asian Journal of Research in Computer Science*, 17(6), 264-292.
37. Omoseebi, A., Ola, G., & Tyler, J. (2025). Rule-Based Systems in AML. Nuskaityta iš https://www.researchgate.net/publication/389869404_Rule-Based_Systems_in_AML
38. Oztas, B., Cetinkaya, D., Adedoyin, F., Budka, M., Aksu, G., & Dogan, H. (2024). Transaction monitoring in anti-money laundering: A qualitative analysis and points of view from industry. *Future Generation Computer Systems*, 159, 161-171.
39. Oztas, B., Cetinkaya, D., Adedoyin, F., Budka, M., Dogan, H., & Aksu, G. (2023). Perspectives from experts on developing transaction monitoring methods for anti-money laundering. *IEEE International Conference on e-Business Engineering (ICEBE)*, (p. 39-46).
40. Pirmorad, E. (2025). Exploring the In-Context Learning Capabilities of LLMs for Money Laundering Detection in Financial Graphs.

41. Sahoo, P., Singh, A. K., Saha, S., Jain, V., Samrat, M., & Chadha, A. (2024). A Systematic Survey of Prompt Engineering in Large Language Models: Techniques and Applications.
42. Sarno, R., Dewandono, R. D., Ahmad, T., Naufal, M. F., & Sinaga, F. (2015). Hybrid Association Rule Learning and Process Mining for Fraud Detection. *IAENG International Journal of Computer Science*, 13(5).
43. Schneider, F., & Windischbauer, U. (2008). Money laundering: some facts. *European Journal of Law and Economics*, 26.
44. Silva, M. C., Tavares, G. M., Gritti, M. C., Ceravolo, P., & Junior, S. B. (2023). Using Process Mining to Reduce Fraud in Digital Onboarding. *Fintech*, 2(1), 120-137.
45. Stephan, S., Lahann, J., & Fettke, P. (2021). A Case Study on the Application of Process Mining in Combination with Journal Entry Test for Financial Auditing. (p. 5718-5727). 54th Hawaii International Conference on System Sciences.
46. Tang, T., Yao, J., & Wang, Y. (2024). Application of Deep Generative Models for Anomaly Detection in Complex Financial Transactions. *4th International Conference on Artificial Intelligence, Internet and Digital Economy (ICAID)* (p. 133-137). IEEE.
47. United Nations Office on Drugs and Crime. (2025). *Money Laundering. Overview*. Nuskaityta iš <https://www.unodc.org/unodc/en/money-laundering/overview.html>
48. Vorobyev, I., & Krivitskaya, A. (2022). Reducing false positives in bank anti-fraud systems based on rule induction in distributed tree-based models. *Computers & Security*, 120.
49. Weber, P., Bordbar, B., & Tino, P. (2013). A principled approach to mining from noisy logs using Heuristics Miner. *Symposium on Computational Intelligence and Data Mining (CIDM)* (p. 119-126). IEEE.
50. Weerapong, S., Porouhan, P., & Premchaiswadi, W. (2012). Process mining using α -algorithm as a tool (A case study of student registration). In *2012 tenth international conference on ICT and Knowledge engineering* (p. 213-220). IEEE.
51. Zhang, C., Qiang, Z., Li, K., Nuthalapati, S. V., Zhang, B., Liu, J., . . . Fan, X. (2025). GEM: Empowering LLM for both Embedding Generation and Language Understanding.

Nuoroda į viešai prieinamą GitHub programinio kodo valdymo saugyklą (angl. repository), kurioje pateikiami programiniai kodai, naudoti procesų gavybos savybių išvedimui, mokėjimų duomenų su tam tikromis užduotimis pateikimui dideliame kalbos modeliui (gpt-4.1-mini) per API sąsają ir finansinių duomenų konvertavimo į JSON formatą programinis kodas:

<https://github.com/ITMartynas/Magistro-Darbas.git>

3 PRIEDAS

Case Id	# of Activities	Throughput Time	First Activity	First Activity Timestamp	Last Activity	Last Activity Timestamp		Case details: 676369110710_Q3	
676369110710_Q3	383		91 d home	7/01/20 15:36:59	misc_pos	9/30/20 22:10:36		Search	Q
676314217768_Q3	95		90 d food_dining	7/02/20 19:23:54	shopping_net	9/30/20 15:08:10		Activities	383 items
676281772837_Q4	347		92 d misc_net	10/01/20 03:53:35	kids_pets	12/31/20 20:43:26		gas_transport	+3M
676275912597_Q4	113		92 d misc_pos	10/01/20 01:15:29	personal_care	12/31/20 17:49:48		home	+3M
676173792455_Q3	279		91 d kids_pets	7/01/20 13:54:16	grocery_net	9/30/20 09:00:09		travel	+3M
675961917837_Q3	95		89 d shopping_pos	7/02/20 01:00:31	misc_net	9/29/20 00:10:11		shopping_net	+3M
6597248210463122_Q4	139		91 d health_fitness	10/01/20 12:32:32	grocery_pos	12/31/20 10:53:52		grocery_net	+3M
6596735789587928_Q3	189		91 d shopping_pos	7/01/20 09:33:44	personal_care	9/30/20 15:21:58		grocery_pos	+3M
6595970453799027_Q3	306		91 d misc_net	7/01/20 10:46:10	travel	9/30/20 13:07:46		kids_pets	+3M
6593939509150446_Q3	108		89 d shopping_net	7/01/20 22:53:51	shopping_net	9/28/20 23:19:43		misc_pos	+3M
6564459919350820_Q2	10		9 d misc_net_fraud	6/21/20 22:32:22	home	6/30/20 13:40:56		shopping_pos	+3M
6544734391390261_Q2	33		9 d home	6/21/20 15:37:14	misc_net	6/30/20 18:53:21		gas_transport_fraud	+3M
6538441737335434_Q4	741		91 d personal_care	10/01/20 12:55:03	personal_care	12/31/20 23:00:37		grocery_pos_fraud	+3M
639046421587_Q4	250		90 d health_fitness	10/02/20 20:53:38	misc_pos	12/31/20 10:31:18		misc_pos_fraud	+3M
630469040731_Q4	446		91 d shopping_pos	10/02/20 05:55:29	health_fitness	12/31/20 22:20:46		misc_net_fraud	+3M
630423337322_Q4	679		92 d misc_net	10/01/20 01:21:15	home	12/31/20 23:17:42			
60495593109_Q3	106		90 d grocery_pos	7/01/20 01:42:18	shopping_net_fraud	9/28/20 23:03:00			
60487002085_Q3	113		87 d shopping_net	7/05/20 20:55:31	kids_pets	9/30/20 23:37:48			
6011975266774121_Q3	104		87 d gas_transport	7/02/20 11:51:49	travel	9/27/20 12:49:58			
6011893664860915_Q3	577		91 d personal_care_fraud	7/01/20 22:51:21	shopping_net	9/30/20 23:56:26			
6011492816282597_Q4	322		92 d misc_net	10/01/20 02:13:32	shopping_net	12/31/20 22:34:58			
6011477612335392_Q3	99		90 d kids_pets	7/01/20 19:02:26	grocery_pos_fraud	9/29/20 23:23:29			
6011399591920186_Q3	390		91 d shopping_pos	7/01/20 14:39:06	home	9/30/20 21:34:32			
6011104316292105_Q3	470		91 d personal_care	7/01/20 17:57:39	home	9/30/20 14:32:10			

H₂O FLOW Flow Cell Data Model Score Admin Help

Untitled Flow

Escape Character 0

Column Headers ☐ Auto

☒ First row contains column names

☐ First row contains data

Options ☐ Enable single quotes as a field quotation character

☒ Delete on done

EDIT COLUMN NAMES AND TYPES

Search by column name...

16	cityPopula	Numeric	88735	53	6841	1745	606
17	job	Enum	Commercial horticulturist	Sports administrator	Therapist, sports	Engineer, automotive	Surveyor, land/geomatics
18	dateOfBirth	Time	1988-04-09	1972-10-05	1999-06-06	1973-12-26	1988-09-06
19	transNum	String	ea11379e8a1b08d584149d65770faee	4a0819d758611bad73aaa547090f3df2	8ec457ccddf3114213845f9d8ea1759c	93ee6c390ef6df5937dab28ec56ccbec	e565ae3e9a18369741363045
20	unixTime	Numeric	1371817121	1372002346	1371818155	1371821560	1371828024
21	merchantLa	String	29,83155	43,294894	36,192355	41,407162	36,815628
22	merchantLo	String	-80,926829	-92,198504	-85,869963	-95,817404	-105,649218
23	isFraud	Enum		0	0	0	0
24	isNightTim	Enum		0	0	0	0
25	timeOfTheD	Enum	fternoon	Afternoon	Afternoon	Afternoon	Afternoon
26	ageCategor	Enum	oung Adult	Middle Age	Young Adult	Middle Age	Young Adult
27	paymentFre	Enum		1	1	1	1
28	fromNetToPi	Enum		0	0	0	0
29	fromPOSToN	Enum		0	0	0	0
30	repeatedCa	Numeric		0	0	0	0

Previous page

Parse

Enum

Time

UUID

String

Invalid

Ready

H2O FLOW

Flow ▾Cell ▾Data ▾Model ▾Score ▾Admin ▾Help ▾

Untitled Flow

📄 🗂️ + ⬆️ ⬇️ ⬅️ 🔍 📄 📁 🗑️ ⏮️ ⏪ ⏩ ⏭ ⚙️

⬅️ Previous 20 Columns ➡️ Next 20 Columns

▶ CHUNK COMPRESSION SUMMARY

▶ FRAME DISTRIBUTION SUMMARY

```
assist splitFrame, "Key_Frame__GALUTINISFEATURESDATASET.hex"
```

✖ Split Frame

Frame: Key_Frame__GALUTINISFEATURESDATASET.hex ▾

Splits:

Ratio	Key
0.7	TRAINING_Key_Frame__GALUTINISFEATURES! ✕
0.30	VALIDATION_Key_Frame__GALUTINISFEATUR

Add new splits

Seed: 416086

Create

```
splitFrame "Key_Frame__GALUTINISFEATURESDATASET.hex", [0.7],  
["TRAINING_Key_Frame__GALUTINISFEATURESDATASET.hex_0.70","VALIDATION_Key_Frame__GALUTINISFEATURESDATASET.hex_0.30"], 416086
```

Split Frames

Type Key Ratio

TRAINING_Key_Frame__GALUTINISFEATURESDATASET.hex_0.700.7

VALIDATION_Key_Frame__GALUTINISFEATURESDATASET.hex_0.300.300000000000000004

Ready

6 PRIEDAS

Flow

Cell

Data

Model

Score

Admin

Help

Untitled Flow

trainings_frame

TRAINING_Key_Frame

GALUTINISFEATURESDATASET.hex_0.70

Id of the training data frame.

validation_frame

VALIDATION_Key_Frame__GALUTINISFEATURESDATASET.hex_0.30

Id of the validation data frame.

n_folds

0

Number of folds for K-fold cross-validation (0 to disable or >= 2).

response_column

isFraud

Response variable column.

ignored_columns

Search...

Names of columns to ignore for training.

Showing page 2 of 5 - 28 ignored.

☒

city

ENUM(849)

☐

state

ENUM(50)

☒

zipCode

INT

☒

customerLatitude

ENUM(910)

☒

customerLongitude

ENUM(910)

☐

cityPopulation

INT

☐

job

ENUM(478)

☐

dateOfBirth

TIME

☒

transNum

STRING

☒

unixTime

INT

☒ All

☐ None

← Previous 10

→ Next 10

Only show columns with more than 0 % missing values.

ignore_const_cols

☒

Ignore constant columns.

n_trees

50

Number of trees.

max_depth

7

Maximum tree depth (0 for unlimited).

min_rows

20

Fewest allowed (weighted) observations in a leaf.

n_bins

20

For numerical columns (real/int), build a histogram of (at least) this many bins, then split at the best point

seed

-1

Seed for pseudo random number generator (if applicable)

H₂O FLOW Flow Cell Data Model Score Admin Help

Untitled Flow

categorical_encoding	AUTO	Encoding scheme for categorical features	☐
custom_metric_func		Reference to custom evaluation function, format: 'language:keyName=funcName'	☐
custom_distribution_func		Reference to custom distribution, format: 'language:keyName=funcName'	☐
export_checkpoints_dir		Automatically export generated models to this directory.	☐
monotone_constraints	Choose... ↔	A mapping representing monotonic constraints. Use +1 to enforce an increasing constraint and -1 to specify a decreasing constraint.	☐
gainslift_bins	-1	Gains/Lift table number of bins. 0 means disabled.. Default value -1 means automatic binning.	☐
auc_type	AUTO	Set default multinomial AUC type.	☐
EXPERT			
class_sampling_factors	1, 5	Desired over/under-sampling ratios per class (in lexicographic order). If not specified, sampling factors will be automatically computed to obtain class balance during training. Requires balance_classes.	☐
max_after_balance_size	5	Maximum relative size of the training data after balancing class counts (can be less than 1.0). Requires balance_classes.	☐
build_tree_one_node	☐	Run on one node only; no network overhead but fewer cpus used. Suitable for small datasets.	☐
sample_rate_per_class		A list of row sample rates per class (relative fraction for each class, from 0.0 to 1.0), for each tree	☐
col_sample_rate_change_per_level	1	Relative change of the column sampling rate for every level (must be > 0.0 and <= 2.0)	☐
max_abs_leafnode_pred	1.7976931348623157e+308	Maximum absolute value of a leaf node prediction	☐
pred_noise_bandwidth	0	Bandwidth (sigma) of Gaussian multiplicative noise ~N(1,sigma) for tree node predictions	☐
calibrate_model	☐	Use Platt Scaling (default) or Isotonic Regression to calculate calibrated class probabilities. Calibration can provide more accurate estimates of class probabilities.	☐
calibration_frame	(Choose...)	Data for model calibration	☐
calibration_method	AUTO	Calibration method to use	☐
in_training_checkpoints_dir		Create checkpoints into defined directory while training process is still running. In case of cluster shutdown, this checkpoint can be used to restart training.	☐
in_training_checkpoints_tree_interval	1	Checkpoint the model after every so many trees. Parameter is used only when in_training_checkpoints_dir is defined	☐
check_constant_response	☒	Check if response column is constant. If enabled, then an exception is thrown if the response column is a constant value. If disabled, then model will train regardless of the response column being a constant value or not.	☐
auto_rebalance	☒	Allow automatic rebalancing of training and validation datasets	☐

Ready

H₂O FLOW Flow ▾ Cell ▾ Data ▾ Model ▾ Score ▾ Admin ▾ Help ▾

Untitled Flow

model_id	gbm-aabe7018-5f30-4456-99f9-e8656d114c82	Destination id for this model; auto-generated if not specified.																				
training_frame	TRAINING_Key_Frame__GALUTINISFEATURESDATASET.hex_0.70 ▾	Id of the training data frame.																				
validation_frame	VALIDATION_Key_Frame__GALUTINISFEATURESDATASET.hex_0.30 ▾	Id of the validation data frame.																				
n_folds	0	Number of folds for K-fold cross-validation (0 to disable or >= 2).																				
response_column	isFraud ▾	Response variable column.																				
ignored_columns	Search... Showing page 2 of 5 - 28 ignored. <table border="1"> <tbody> <tr><td>☒ city</td><td>ENUM(849)</td></tr> <tr><td>☐ state</td><td>ENUM(50)</td></tr> <tr><td>☒ zipCode</td><td>INT</td></tr> <tr><td>☒ customerLatitude</td><td>ENUM(910)</td></tr> <tr><td>☒ customerLongitude</td><td>ENUM(910)</td></tr> <tr><td>☐ cityPopulation</td><td>INT</td></tr> <tr><td>☐ job</td><td>ENUM(478)</td></tr> <tr><td>☐ dateOfBirth</td><td>TIME</td></tr> <tr><td>☒ transNum</td><td>STRING</td></tr> <tr><td>☒ unixTime</td><td>INT</td></tr> </tbody> </table> ☒ All ☐ None ← Previous 10 Next 10 →		☒ city	ENUM(849)	☐ state	ENUM(50)	☒ zipCode	INT	☒ customerLatitude	ENUM(910)	☒ customerLongitude	ENUM(910)	☐ cityPopulation	INT	☐ job	ENUM(478)	☐ dateOfBirth	TIME	☒ transNum	STRING	☒ unixTime	INT
☒ city	ENUM(849)																					
☐ state	ENUM(50)																					
☒ zipCode	INT																					
☒ customerLatitude	ENUM(910)																					
☒ customerLongitude	ENUM(910)																					
☐ cityPopulation	INT																					
☐ job	ENUM(478)																					
☐ dateOfBirth	TIME																					
☒ transNum	STRING																					
☒ unixTime	INT																					
ignore_const_cols	☒	Ignore constant columns.																				
ntrees	50	Number of trees.																				
max_depth	7	Maximum tree depth (0 for unlimited).																				
min_rows	20	Fewest allowed (weighted) observations in a leaf.																				
nbins	20	For numerical columns (real/int), build a histogram of (at least) this many bins, then split at the best point																				
seed	-1	Seed for pseudo random number generator (if applicable)																				
learn_rate	0.05	Learning rate (from 0.0 to 1.0)																				