

Avtoenkoder Əsaslı Hibrid Yanaşma ilə Kredit Kartı Fırıldaqçılığının Aşkarlanması Metodu

Eltun Əhmədov

Informasiya Texnologiyaları İnstitutu, Bakı, Azərbaycan
eltunehmedov95@gmail.com

Xülasə— Məqalədə kredit kartı fırıldaqçılığının aşkarlanması üçün avtokodlayıcıya əsaslanan hibrid model təklif olunur. Təklif olunan avtokodlayıcı arxitekturası ilkin mərhələdə verilənlərdən informativ xüsusiyyətlərin çıxarılmasını və ölçünün azaldılmasını təmin edir və əldə olunan xüsusiyyətlər çoxlaylı perseptron və ekstremal qradient artırma maşın öyrənməsi təsnifat modellərinə ötürülərək fırıldaq əməliyyatlarının aşkarlanması üçün istifadə olunur. Eksperimental nəticələr göstərir ki, avtokodlayıcıya əsaslanan hibrid model ənənəvi metodlarla müqayisədə F1-ölçü göstəricisi üzrə daha yüksək nəticə göstərdiyini nümayiş etdirir. Bundan əlavə, təklif olunan metod verilənlərin ölçüsünü azaltmaqla hesablama mürəkkəbliyinin optimallaşdırılmasını təmin edir.

Açar sözlər— kredit kartı fırıldaqçılığı; maşın öyrənməsi; avtokodlayıcı; XGBoost.

I. GİRİŞ

Son illərdə rəqəmsal ödəniş sistemlərinin və onlayn bankçılıq xidmətlərinin sürətli inkişafı kredit kartı əməliyyatlarının həcmnin əhəmiyyətli dərəcədə artmasına səbəb olmuşdur. Bu inkişaf maliyyə əməliyyatlarında əlçatanlığı artırmaqla yanaşı, kredit kartı fırıldaqçılığı hallarının artmasına və maliyyə sistemlərinin kibertəhlükəsizlik baxımından daha həssas vəziyyətə gəlməsinə səbəb olmuşdur. Bu cür fırıldaqçılıq hadisələri maliyyə institutları üçün ciddi iqtisadi itkilərlə nəticələnir və istifadəçi etibarının zəifləməsinə səbəb olur. Fırıldaq əməliyyatları çox vaxt real istifadəçi davranışını təqlid etdiyindən, onların avtomatik aşkarlanması mürəkkəbləşir və bu proses yüksək dəqiqlikli süni intellekt modellərinin inkişafını tələb edir.

Kredit kartı fırıldaqçılığının aşkarlanması problemi elmi ədəbiyyatda adətən ikili təsnifat (binary classification) çərçivəsində formalaşdırılır. Bu yanaşmada hər bir tranzaksiya “normal” və “fırıldaq” siniflərindən birinə aid edilir. Bununla belə, problemin əsas metodoloji çətinliyi verilənlər dəstinin ciddi şəkildə disbalans xarakter daşmasıdır. Belə ki, fırıldaq əməliyyatları ümumi tranzaksiya həcmi ilə müqayisədə olduqca kiçik paya malik olur. Bu disbalans klassik maşın öyrənməsi alqoritmlərinin ümumiləşdirmə qabiliyyətini məhdudlaşdırır və modellərin dominant sinifə (majority class) meyllənməsinə səbəb olur. Nəticədə nadir sinifin (minority class) aşkarlama dəqiqliyi əhəmiyyətli dərəcədə azalır. Digər mühüm problem verilənlərin çoxölçülü olmasıdır. Hər bir tranzaksiyanın çoxsaylı xüsusiyyətlər ilə təsvir edilməsi çoxölçülü verilənlər fəzasının yaranmasına səbəb olur və bu, “lənətə gəlmiş ölçü” (curse of dimensionality) kimi tanınan problemi ortaya çıxarır. Ölçülərin

artması ilə verilənlər fəzasında nümunələr seyrəkləşir, bu da siniflər arasındakı sərhədlərin dəqiq müəyyən edilməsini çətinləşdirir və modelin ümumiləşdirmə qabiliyyətinə mənfi təsir göstərir. Nəticədə modelin həddindən artıq öyrənmə (overfitting) riski artır. Bundan əlavə, çoxölçülü verilənlər daha çox hesablama resursu tələb edir və təlim prosesinin müddətini uzadır.

Bu problemlərin həlli istiqamətində son illərdə dərin öyrənmə (deep learning) əsaslı metodlara, xüsusilə avtokodlayıcı arxitekturalarına qarşı artan elmi maraq müşahidə olunur. AE-lər giriş verilənlərinin aşağıölçülü və sıxlaşdırılmış latent (gizli) təsvirlərini öyrənməklə, həm xüsusiyyətlərin avtomatik çıxarılması (feature extraction), həm də rekonstruksiya əsasında küyün azaldılması (denoising) funksiyalarını yerinə yetirən neyron şəbəkə modelləridir.

Bu yanaşma çoxölçülü və qeyri-xətti struktura malik maliyyə tranzaksiya verilənlərinin effektiv modelləşdirilməsini təmin edir. Xüsusilə, AE-lər mürəkkəb paylanmaların gizli layda nümunələrini öyrənmə qabiliyyəti sayəsində, ənənəvi xətti ölçü azaldılması metodları ilə müqayisədə daha zəngin və diskriminativ xüsusiyyətlərin əldə edilməsinə imkan verir. Bununla yanaşı, rekonstruksiya xətası əsasında qurulan anomaliya aşkarlama mexanizmi, fırıldaq əməliyyatlarının müəyyən edilməsində effektiv alternativ yanaşma kimi çıxış edir.

AE iki əsas hissədən ibarət olan neyron şəbəkə arxitekturasıdır: kodlayıcı (encoder) və dekodlayıcı (decoder). Kodlayıcının əsas vəzifəsi giriş verilənlərini aşağıölçülü gizli laya çevirməkdir. Dekodlayıcı isə bu gizli layı qəbul edərək orijinal giriş verilənlərini mümkün qədər dəqiq şəkildə bərpa etməyə çalışır. Beləliklə, modelin məqsədi giriş və çıxış arasındakı fərqi minimallaşdırmaqdır. Bu fərqi azaldılması modelin verilənləri daha effektiv öyrənməsini təmin edir və gizli layın təmsil gücünü artırır.

AE arxitekturasında vacib anlayışlardan biri “darboğaz” (bottleneck) layıdır. Bu lay kodlayıcı və dekodlayıcı layları arasında yerləşərək verilənlərin aşağıölçülü gizli fəzada sıxlaşdırılmış şəkildə təmsil olunmasını təmin edir. “Darboğaz” mexanizmi modelin giriş verilənlərini tam şəkildə yadda saxlamaq əvəzinə yalnız ən informativ xüsusiyyətləri saxlayır. Bu yanaşma xüsusiyyətlərin çıxarılması və ölçü azaldılması üçün effektiv mexanizm formalaşdırır.

Bu tədqiqatda AE əsaslı yanaşma kredit kartı fırıldaqçılığının aşkarlanması üçün tətbiq olunur. Tədqiqatın məqsədi verilənlərdən daha aşağıölçülü xüsusiyyətlərin

çıxarılması və bu xüsusiyyətlər əsasında kredit kartı fırıldaqçılığının daha dəqiq şəkildə aşkarlanmasıdır.

II. ƏLAQƏLİ İŞLƏR

[1]-də fırıldaqçılıq əməliyyatlarının aşkarlanması üçün iki mərhələli model təklif olunur. Birinci mərhələdə AE vasitəsilə çoxölçülü xüsusiyyətlər aşağıölçülü gizli fəzaya transformasiya edilir, ikinci mərhələdə isə əldə olunan xüsusiyyət vektoru MLP, K-Nearest Neighbors və Logistic Regression kimi təsnifat alqoritmlərinə daxil edilərək qərarvermə prosesi həyata keçirilir.

[2]-də kredit kartı fırıldaqçılığının aşkarlanması üçün AE və probablistik LightGBM əsaslı AED-LGB modeli təklif olunur. Eyni qaydada bu yanaşmada AE çoxölçülü verilənlərdən aşağıölçülü xüsusiyyətlər çıxarır və əldə olunan xüsusiyyətlər LightGBM alqoritmı ilə təsnif edilir. [3]-də Bal-VAE-Attention adlı yeni yanaşma təklif olunur. Bu yanaşma disbalans verilənlərdə paylanma xüsusiyyətlərini nəzərə alan variasional avtokodlayıcı (variational autoencoder) əsaslı nümunə artırma (oversampling), avtomatik xüsusiyyət seçimi və çoxbaşı diqqət mexanizmi (multi-head attention) əsasında fırıldaq əməliyyatlarının mürəkkəb daxili strukturlarını öyrənməyə imkan yaradır.

[4]-də fırıldaqçılıq əməliyyatlarının aşkarlanması üçün iki mərhələli dərin öyrənmə modeli təklif olunur. Birinci mərhələdə dərin AE vasitəsilə verilənlər aşağıölçülü xüsusiyyət fəzasına çevrilir, ikinci mərhələdə isə bu transformasiya olunmuş xüsusiyyətlər üzərində dərin öyrənmə əsaslı klassifikatorlar tətbiq edilir. Nəticələr göstərir ki, AE əsaslı model həm orijinal verilənlər, həm də əsas komponentlər analizi (Principal Component Analysis, PCA) əsaslı yanaşma ilə müqayisədə bütün performans göstəricilərində daha yüksək təsnifat effektivliyi təmin edir.

[5]-də kredit kartı fırıldaqçılığının aşkarlanması üçün maşın öyrənməsinə əsaslanan əks əlaqə mexanizminə malik yanaşma təklif olunur və müxtəlif təsnifat alqoritmlərinin (Random Forest, Support Vector Machine, Artificial Neural Network və s.) performansını müqayisəli şəkildə qiymətləndirilir. Təklif olunan modeldə əvvəlki təsnifat nəticələri sistemə geri ötürülərək öyrənmə prosesi adaptiv şəkildə yenilənir, bu isə xüsusilə disbalans verilənlər üzərində aşkarlama dəqiqliyini və ümumi təsnifat performansını artırır.

III. METODOLOGİYA

İlkin mərhələdə verilənlər giriş xüsusiyyətləri (input features) və sinif (label) üzrə ayrılır. Daha sonra verilənlər toplusu 80% təlim və 20% test altqoxluqlarına bölünür. Təlim altqoxluğu isə əlavə olaraq 80% təlim və 20% validasiya hissələrinə ayrılır. Nəticədə ümumi bölgü 64% təlim, 16% validasiya və 20% test şəklində formalaşır. Bu yanaşma sinif paylanmasının bütün altqoxluqlarda qorunmasını təmin edərək modelin qiymətləndirilməsinin daha etibarlı və real şəraitə uyğun aparılmasına imkan verir. Daha sonra giriş dəyişənləri üçün “MinMaxScaler” əsaslı normallaşdırma tətbiq edilir. Normallaşdırma parametrləri yalnız təlim verilənləri əsasında müəyyən edilir və eyni parametrlər validasiya və test verilənlərinə tətbiq olunur. Bu prosedur məlumat sızmasının (data leakage) qarşısını almaq üçün standart metodoloji yanaşma

hesab olunur.

Növbəti mərhələdə AE arxitekturası qurularaq təlim edilir. Təlim prosesi tamamlandıqdan sonra modelin yalnız kodlayıcı hissəsi çıxarılaraq sabit xüsusiyyət çıxarıcı (feature extractor) kimi istifadə edilir.

Əldə olunan gizli xüsusiyyətlər növbəti mərhələdə kredit kartı fırıldaqçılığının aşkarlanması üçün iki maşın öyrənməsi alqoritmına giriş kimi təqdim edilir: MLP [6] və XGBoost [7].

IV. EKSPERİMENT

A. Verilənlər toplusu

Bu tədqiqatda istifadə edilən verilənlər toplusu Kaggle platformasında açıq şəkildə mövcuddur və Avropadakı kart sahibləri tərəfindən 2013-cü ilin sentyabr ayında iki gün ərzində həyata keçirilmiş kredit kartı əməliyyatlarını əhatə edir [8]. Verilənlər toplusu 284807 əməliyyatı əhatə edir və onların yalnız 492-si fırıldaq əməliyyatı kimi qeydə alınmışdır. Bu göstərici verilənlərdə ciddi sinif disbalansının mövcud olduğunu göstərir, belə ki, fırıldaq əməliyyatları ümumi əməliyyatların cəmi 0.173%-ni təşkil edir. Bu isə nadir sinifin düzgün öyrənilməsinə çətinləşdirən mühüm amillərdən biridir.

Verilənlər toplusu 31 ədədi atributdan ibarətdir. Atributların böyük əksəriyyəti həssas məlumatların məxfiliyinin qorunması məqsədilə PCA vasitəsilə transformasiya olunmuşdur. Yalnız “Time”, “Amount” və “Class” xüsusiyyətləri orijinal formada saxlanılmışdır.

B. Modelin quruluşu

Verilənlərin ilkin emalı mərhələsində istifadə olunan AE modelinin hiperparametrləri və onların qiymət diapazonları Cədvəl 1-də təqdim olunmuşdur.

CƏDVƏL I. AE MODELİ ÜÇÜN HIPERPARAMETRLƏRİN AXTARIŞ DİAPAZONLARI

Hiperparametr	Axtarış Diapazonu
Kodlayıcı və dekodlayıcıda layların sayı	{2,3,4}
Hər layda neyronların sayı	[1-512]
“Darboğaz” layda neyronların sayı	[15-22]
“Darboğaz” layının aktivasiya funksiyası	{Sigmoid, Linear}
Digər layların aktivasiya funksiyası	{Sigmoid, ReLU, Tanh}
Paket ölçüsü (Batch size)	[32-512]
Epoxa sayı (epochs)	[30-100]
İtki funksiyası (loss function)	{MSE, MAE, BCE}

Hiperparametrlərin tənzimlənməsi mərhələsində itki funksiyası kimi orta kvadratik xəta (Mean Squared Error – MSE), orta mütləq xəta (Mean Absolute Error – MAE) və binar çarpaz entropiya (Binary Cross Entropy – BCE) funksiyaları nəzərdən keçirilmişdir. Bu və digər hiperparametrlər modelin öyrənmə qabiliyyətinə, sıxlaşdırma səviyyəsinə və nəticədə əldə olunan gizli fəzanın keyfiyyətinə birbaşa təsir göstərir. Seçilmiş hiperparametrlər əsasında model ilkin verilənlər üzərində təlim edilmiş və nəticədə verilənlərin daha kompakt və informativ

gizli fəzası əldə olunmuşdur. Təlimdən sonra modelin “darboğaz” layının çıxışı istifadə edilərək giriş verilənləri daha aşağıölçülü fəzaya inikas edilmişdir.

Hiperparametrlərin tənzimlənməsi nəticəsində əldə edilmiş optimal model arxitekturasına əsasən, AE modeli kodlayıcı, “darboğaz” və dekodlayıcı olmaqla üç hissədən ibarətdir. Kodlayıcı hissəsi girişdən gizli fəzaya qədər 512, 128 və 64 neyronlu üç gizli laydan, dekodlayıcı hissəsi isə gizli fəzadan giriş ölçüsünün bərpasına qədər simmetrik şəkildə 64, 128 və 512 neyronlu üç gizli laydan təşkil olunmuşdur. Modeldə “darboğaz” layının ölçüsü 18 neyron olaraq müəyyən edilmişdir.

Modelin təlimi zamanı optimallaşdırma alqoritmi kimi Adam, itki funksiyası kimi MAE funksiyasından istifadə edilmişdir. Təlim prosesi 50 epoxa ərzində aparılmış və paket ölçüsü 256 olaraq seçilmişdir. Ən yaxşı nəticə göstərən AE modelinin arxitekturası Cədvəl 2-də göstərilmişdir.

CƏDVƏL II. AE MODELİNİN TOPOLOGİYASI

Təklif olunan AE arxitekturası		
Laylar	Neyron sayı	Aktivasiya funksiyası
Giriş	30	-
Kodlayıcı	512	Tanh
Kodlayıcı	128	Tanh
Kodlayıcı	64	Tanh
“Darboğaz”	18	Sigmoid
Dekodlayıcı	64	Tanh
Dekodlayıcı	128	Tanh
Dekodlayıcı	512	Tanh
Çıxış	30	Sigmoid

C. Nəticələrin qiymətləndirilməsi

Tədqiqat işində aparılan eksperimentlər Python proqramlaşdırma mühitində həyata keçirilmişdir. Verilənlərin emalı, modellərin qurulması və qiymətləndirilməsi mərhələlərində Pandas, NumPy, Scikit-learn, TensorFlow və XGBoost kitabxanalarından istifadə olunmuşdur. Əldə olunan nəticələrin qiymətləndirilməsi məqsədilə dəqiqlik (Accuracy), həssaslıq (Precision), tamlıq (Recall) və F1-ölçü (F1-score) metrikalarından istifadə edilmişdir [9].

Cədvəl 3-ə əsasən, təklif olunan AE-MLP modeli həssaslıq (0.8119) və tamlıq (0.8367) göstəriciləri arasında daha az fərq göstərərək daha balanslı performans təmin edir. Bu isə müsbət siniflərin daha etibarlı və sabit şəkildə aşkarlanmasını təmin edir.

F1-ölçü baxımından AE-MLP modeli 0.8241 nəticəsi ilə həm baza MLP (0.7882), həm də PCA-MLP (0.8125) modellərini üstələyir. Bu isə AE əsaslı yanaşmanın xüsusiyyətlərin daha informativ şəkildə kodlaşdırılması hesabına həssaslıq və tamlıq arasında daha optimal kompromis yaratdığını göstərir. Xüsusilə tamlıq göstəricisinin yüksək olması (0.8367) modelin müsbət sinfə aid nümunələri (fırıldaqçılıq halları) daha effektiv şəkildə aşkarladığını təsdiqləyir.

Ümumilikdə, AE-MLP modeli PCA metoduna əsaslanan yanaşma ilə müqayisədə daha stabil və balanslı performans nümayiş etdirir və F1-ölçü üzrə əldə olunan üstünlük təklif olunan metodun effektivliyini sübut edir. Bu nəticələr AE əsaslı xüsusiyyət çıxarılmasının MLP modelinin ümumiləşdirmə qabiliyyətini artırdığını və fırıldaqçılığın aşkarlanma keyfiyyətini daha optimal səviyyəyə çatdırdığını göstərir.

CƏDVƏL III. FIRILDAQÇILIĞIN AŞKARLANMASINDA MLP MODELİ

Ölçü sayı	Model	Dəqiqlik	Həssaslıq	Tamlıq	F1-ölçü
30	(baza) MLP	0.9992	0.7619	0.8163	0.7882
18	PCA-MLP	0.9994	0.8298	0.7959	0.8125
18	AE-MLP	0.9994	0.8119	0.8367	0.8241

Cədvəl 4-ə əsasən, AE-XGBoost modeli həssaslıq (0.8936) və tamlıq (0.8571) göstəriciləri baxımından digər modellərlə müqayisədə daha balanslı performans göstərir. Xüsusilə PCA-XGBoost modeli ilə müqayisədə tamlıq göstəricisində müşahidə olunan artım (0.7857→0.8571) müsbət siniflərin daha effektiv şəkildə aşkarlanmasını təmin edir.

F1-ölçü baxımından AE-XGBoost modeli 0.8750 nəticəsi ilə həm baza XGBoost (0.8586), həm də PCA-XGBoost (0.8462) modellərini üstələyir. Bu isə AE əsaslı yanaşmanın həssaslıq və tamlıq arasında daha optimal tarazlıq yaratdığını və modelin ümumi təsnifat performansını yaxşılaşdırdığını göstərir.

Ümumilikdə, nəticələr AE-XGBoost modelinin xüsusiyyətlərin daha informativ şəkildə çıxarılması hesabına daha stabil və effektiv performans nümayiş etdirdiyini təsdiq edir.

CƏDVƏL IV. FIRILDAQÇILIĞIN AŞKARLANMASINDA XGBOOST MODELİ

Ölçü sayı	Model	Dəqiqlik	Həssaslıq	Tamlıq	F1-ölçü
30	(baza) XGBoost	0.9995	0.8817	0.8367	0.8586
18	PCA-XGBoost	0.9995	0.9167	0.7857	0.8462
18	AE-XGBoost	0.9996	0.8936	0.8571	0.8750

Cədvəl V-də təqdim olunan nəticələrin müqayisəli təhlili göstərir ki, mövcud tədqiqatlarla müqayisədə təklif olunan yanaşma daha yüksək və balanslı performans nümayiş etdirir. Ədəbiyyatda təqdim olunan modellər arasında [1] tərəfindən təklif olunan AE əsaslı MLP modeli F1-ölçü (0.82) və həssaslıq-tamlıq balans baxımından nisbətən stabil nəticələr əldə etmişdir. Eyni zamanda, [10] tərəfindən təklif olunan XGBoost modeli F1-ölçü (0.85) üzrə daha yüksək nəticə göstərmişdir.

Bununla belə, [11] tərəfindən təqdim edilən Random Forest modelində həssaslıq göstəricisinin olduqca aşağı olması (0.09) F1-ölçü göstəricisinin kəskin azalmasına (0.17) səbəb olmuş və modelin disbalans performansını ortaya qoymuşdur.

Mövcud yanaşmalarla müqayisədə, təklif olunan AE-XGBoost modeli həssaslıq (0.8936), tamlıq (0.8571) və F1-ölçü (0.8750) göstəriciləri ilə daha stabil və balanslı performans göstərir. Xüsusilə [10] modeli ilə müqayisədə tamlıq və F1-ölçü üzrə əldə olunan üstünlük, AE əsaslı xüsusiyyət çıxarılmasının

modelin ümumi effektivliyini artırdığını və daha etibarlı təsnifat nəticələri verdiyini təsdiq edir.

CƏDVƏL V. TƏKLİF OLUNAN MODELİN MÖVCUD METODLARLA MÜQAYİSƏSİ

İstinad	İl	Model	Dəqiqlik	Həssaslıq	Tamlıq	F1-ölçü
[1]	2020	MLP	0.99	0.85	0.80	0.82
[10]	2023	XGBoost	0.99	0.89	0.81	0.85
[11]	2023	Random Forest	0.96	0.09	0.97	0.17

NƏTİCƏ

Bu tədqiqat çərçivəsində kredit kartı fırıldaqçılığının aşkarlanması üçün AE əsaslı hibrid yanaşma təklif edilmiş, MLP və XGBoost təsnifat modelləri ilə integrasiya olunaraq qiymətləndirilmişdir. Eksperimental nəticələr göstərir ki, AE vasitəsilə əldə olunan gizli nümunələr modellərin fərqləndirmə qabiliyyətini artırır və təsnifat göstəricilərində daha balanslı nəticələrin əldə olunmasına şərait yaradır.

Bundan əlavə, ölçü azaldılması yanaşmasının tətbiqi çoxölçülü verilənlərin daha kompakt və informativ formada təqdim olunmasını təmin edərək hesablama mürəkkəbliyini azaldır və təlim prosesinin səmərəliliyini artırır. Bu isə həm resurs istifadəsinin optimallaşdırılmasına, həm də praktik tətbiq imkanlarının genişlənməsinə töhfə verir.

Ümumilikdə, əldə olunan nəticələr təklif olunan metodun həm təsnifat dəqiqliyi, həm də hesablama səmərəliliyi baxımından üstün performans nümayiş etdirdiyini təsdiqləyir. Bu isə təklif olunan AE əsaslı modelin kredit kartı fırıldaqçılığının aşkarlanması problemi üçün effektiv və perspektivli bir yanaşma olduğunu göstərir.

ƏDƏBİYYAT

- [1] S. Misra, S. Thakur, M. Ghosh, and S. K. Saha, "An Autoencoder Based Model for Detecting Fraudulent Credit Card Transaction," *Procedia Computer Science*, vol. 167, pp. 254–262, 2020, doi: 10.1016/j.procs.2020.03.219.
- [2] H. Du, L. Lv, A. Guo, and H. Wang, "AutoEncoder and LightGBM for Credit Card Fraud Detection Problems," *Symmetry*, vol. 15, no. 4, p. 870, Apr. 2023, doi: 10.3390/sym15040870.
- [3] S. Shi, W. Luo, and G. Pau, "An attention-based balanced variational autoencoder method for credit card fraud detection," *Applied Soft Computing*, vol. 177, p. 113190, Jun. 2025, doi: 10.1016/j.asoc.2025.113190.
- [4] H. Fanai and H. Abbasimehr, "A novel combined approach based on deep Autoencoder and deep classifiers for credit card fraud detection," *Expert Systems with Applications*, vol. 217, p. 119562, 2023, doi: 10.1016/j.eswa.2023.119562.

- [5] N. K. Trivedi, S. Simaiya, U. K. Lilhore, and S. K. Sharma, "An efficient credit card fraud detection model based on machine learning methods," *International Journal of Advanced Science and Technology*, vol. 29, no. 5, pp. 3414–3424, 2020, ISSN: 2005-4238.
- [6] D. E. Rumelhart, G. E. Hinton, and R. J. Williams, "Learning representations by back-propagating errors," *Nature*, vol. 323, pp. 533–536, 1986.
- [7] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining (KDD)*, 2016, pp. 785–794.
- [8] A. Dal Pozzolo, O. Caelen, R. Johnson, G. Bontempi, and S. Waterschoot, "Credit Card Fraud Detection Dataset," *Kaggle Dataset*, 2015. <https://www.kaggle.com/datasets/mlg-ulb/creditcardfraud>
- [9] D. M. Powers, "Evaluation: From Precision, Recall and F-measure to ROC, informedness, markedness and correlation," *Journal of Machine Learning Technologies*, vol. 2, no. 1, pp. 37–63, 2011.
- [10] C. Do, D. Ngoc, and N. Duy, "A new approach for detecting credit card fraud transaction," *International Journal of Computer Applications*, vol. 14, no. 3, pp. 133–146, Mar. 2023.
- [11] J. K. Afriyie, K. Tawiah, W. A. Pels, S. Addai-Henne, H. A. Dwamena, E. O. Owiredo, S. A. Ayeh, and J. Eshun, "A supervised machine learning algorithm for detecting and predicting fraud in credit card transactions," *Decision Analytics Journal*, vol. 6, p. 100163, 2023, doi: 10.1016/j.dajour.2023.100163.

Autoencoder-Based Hybrid Approach for Credit Card Fraud Detection

Eltun Ahmadov

Institute of Information Technology, Baku, Azerbaijan
eltunehmedov95@gmail.com

Abstract- Credit card fraud detection is a complex classification problem that requires the analysis of multidimensional, imbalanced, and dynamic data. The variable nature of fraudulent transactions necessitates the development of effective and robust models. For this purpose, this study proposes an autoencoder-based hybrid model for credit card fraud detection. The proposed autoencoder architecture enables the extraction of informative features from the data and preforms dimensionality reduction in the initial stage. The resulting features are then fed into MLP and XGBoost machine learning classification models to detect fraudulent transactions. Experimental results show that the autoencoder-based hybrid model achieves higher performance in terms of the F1-score compared to traditional methods. In addition, the proposed method reduces computational complexity by lowering the dimensionality of data.

Keywords- credit card fraud; machine learning; autoencoder; XGBoost.