



**UNIVERSITETI I TIRANËS
FAKULTETI I SHKENCAVE TË NATYRËS
DEPARTAMENTI I INFORMATIKËS**

DISERTACION

TEMA

**“APLIKIMI I TEKNIKAVE TË INTELIGJENCËS ARTIFICIALE
NË NDIHMESE TË REALIZIMIT TË VENDIMARRJES AUTOMATIKE”**

Kandidati

M.Sc.Ilma LILI

Udhëheqës shkencor:

Prof. Asoc. Dr. Endrit XHINA

Tiranë, 2026



UNIVERSITETI I TIRANËS
FAKULTETI I SHKENCAVE TË NATYRËS
DEPARTAMENTI I INFORMATIKËS

Disertacion i paraqitur nga
Znj. Ilma Lili

për marrjen e gradës shkencore

“Doktor i shkencave”

Nënfusha e studimit

“Informatikë”

Tema: “Aplikimi i teknikave të Inteligjencës Artificiale në ndihmesë të realizimit të vendimarrjes automatike”

Mbrohet në datë..... para jurisë:

1. _____ Kryetar
2. _____ Anëtar (Oponent)
3. _____ Anëtar (Oponent)
4. _____ Anëtar
5. _____ Anëtar

Deklarata e origjinalitetit të punimit

Unë, Ilma Lili, deklaroj, nën përgjegjësinë time personale, se punimi i doktoratës me temë: *“Aplikimi i teknikave të Inteligjencës Artificiale në ndihmesë të realizimit të vendimarrjes automatike”*, me udhëheqës shkencor Prof. Asoc. Dr. Endrit Xhina, është përfundim i punës kërkimore shkencore origjinale. Punimi nuk është prezantuar asnjëherë përpara një institucioni tjetër për vlerësim apo të jetë botuar i tëri dhe nuk ka në përmbajtje të tij materiale të shkruara nga autorë të tjerë, përveç rasteve të cilat janë të referuara dhe të cituara.

Ilma LILI

Endrit XHINA

Abstrakt

Ky disertacion trajton rolin e teknikave të Inteligjencës Artificiale për të ndihmuar në realizimin e sistemeve të vendimmarrjes automatike në kontekste të ndryshme aplikative. Studimi fokusohet në analizimin e marrëdhënies ndërmjet performancës së modeleve inteligjente, transparencës së vendimeve dhe integritit të sistemeve autonome në mjedise reale. Në kuadër të punimit janë trajtuar konceptet teorike të vendimmarrjes automatike, modelet probabilistike, teknikat e të Mësuarit të Makinës dhe të Mësuarit të Thelluar, si dhe qasjet e Inteligjencës Artificiale të Shpjegueshme (XAI). Nga ana metodologjike, disertacioni paraqet tre raste studimore kryesore. Rasti i parë trajton integrimin e metodave XAI në chatbot-e inteligjente për sektorin e hotelierisë, me qëllim rritjen e transparencës dhe besueshmërisë së vendimeve automatike. Rasti i dytë fokusohet në analizën e sentimentit nga e folura në gjuhën shqipe, duke përdorur modele të Mësuarit të Thelluar dhe përfaqësime akustike si MFCC dhe HuBERT në një kontekst me burime të kufizuara gjuhësore. Rasti i tretë propozon një framework të bazuar në konceptin e Binjakut Dixhital (Digital Twin) për monitorimin e cilësisë së ajrit në ambiente të brendshme dhe realizimin e vendimmarrjes automatike përmes sensorëve IoT, modeleve parashikuese dhe moduleve adaptive të kontrollit. Rezultatet tregojnë se teknikat e Inteligjencës Artificiale mund të përdoren në mënyrë efektive për automatizimin e proceseve vendimmarrëse, duke rritur saktësinë, shpejtësinë dhe aftësinë parashikuese të sistemeve së bashku me integrimin e metodave të shpjegueshmërisë.

Fjalë kyçe: *Inteligjenca Artificiale, vendimmarrja automatike, Machine Learning, Deep Learning, XAI (eXplainable AI), Digital Twin, analiza e sentimentit, IoT, HuBERT, chatbot inteligjent.*

Abstract

This dissertation examines the role of Artificial Intelligence techniques in supporting the implementation of automated decision-making systems across different application contexts. The study focuses on analyzing the relationship between the performance of intelligent models, decision transparency, and the integration of autonomous systems into real-world environments. Within the scope of this work, the theoretical concepts of automated decision-making, probabilistic models, Machine Learning and Deep Learning techniques, as well as Explainable Artificial Intelligence (XAI) approaches are addressed. From a methodological perspective, the dissertation presents three main case studies. The first case study examines the integration of XAI methods into intelligent chatbots for the hospitality sector, with the aim of increasing the transparency and reliability of automated decisions. The second case study focuses on speech sentiment analysis in the Albanian language by using Deep Learning models and acoustic representations such as MFCC and HuBERT within a low-resource language context. The third case study proposes a framework based on the concept of the Digital Twin for indoor air quality monitoring and the implementation of automated decision-making through IoT sensors, predictive models, and adaptive control modules. The results demonstrate that Artificial Intelligence techniques can be effectively used to automate decision-making processes by improving the accuracy, speed, and predictive capabilities of the systems, together with the integration of explainability methods.

Keywords: *Artificial Intelligence, automated decision-making, Machine Learning, Deep Learning, XAI (eXplainable Artificial Intelligence), Digital Twin, sentiment analysis, IoT, HuBERT, intelligent chatbot.*

Falënderime

Ky disertacion përfaqëson jo vetëm një rrugëtim akademik dhe shkencor, por edhe një rrugëtim personal të mbushur me sfida, sakrifica, përkushtim dhe ëndrra. Asnjë arritje nuk ndërtohet vetëm, ndaj dëshiroj të shpreh mirënjohjen time më të thellë për të gjithë ata njerëz që kanë qenë pranë meje gjatë këtij rrugëtimi.

Së pari, falënderoj Zotin për forcën, shëndetin dhe mundësinë që më dha për të përfunduar këtë rrugëtim akademik.

Dëshiroj të falënderoj me respekt dhe mirënjohje të veçantë udhëheqësin tim shkencor Prof Asoc Dr. Endrit Xhina, për besimin e vazhdueshëm që ka pasur tek unë, për mbështetjen, sugjerimet dhe idetë e vlefshme që më kanë ndihmuar të ec përpara edhe në momentet më të vështira. Besimi dhe inkurajimi i tij kanë qenë një forcë shtytëse e rëndësishme në realizimin e këtij punimi. Një falënderim i veçantë shkon edhe për Prof. Asoc. Dr. Alda Kikën, shefen e Departamentit të Informatikës, për mbështetjen, mirëkuptimin dhe motivimin e vazhdueshëm gjatë këtij procesi akademik.

Një falënderim i veçantë shkon për familjen time, e cila ka qenë gjithmonë mbështetja ime më e fortë gjatë këtij rrugëtimi. Falënderoj me gjithë zemër mamën Sabahete, babin Osman, nanën Leka dhe vëllain tim Ervis për dashurinë, përkushtimin dhe mbështetjen e tyre të vazhdueshme në çdo hap të jetës sime.

Megjithatë, falënderimi më i veçantë i dedikohet bashkëshortit tim, Erionit, i cili me durimin, përkushtimin dhe mbështetjen e tij të vazhdueshme ka qenë gjithmonë pranë meje gjatë këtij rrugëtimi. Inkurajimi, mirëkuptimi dhe besimi i tij kanë përfaqësuar një mbështetje të çmuar, duke më dhënë forcën dhe motivimin për të përballuar sfidat dhe për të vazhduar drejt realizimit të objektivave të mia akademike dhe personale. Prania e tij në jetën time ka qenë një nga shtyllat më të rëndësishme të këtij rrugëtimi.

Një mendim plot dashuri shkon edhe për fëmijët e mi të shtrenjtë, Erginin dhe Skerbin, të cilëve shpeshherë u kam marrë kohë nga fëmijëria e tyre për t'iu përkushtuar këtij punimi. Dashuria dhe buzëqeshja e tyre kanë qenë motivimi im më i madh për të mos u ndalur kurrë.

Falënderoj gjithashtu kolegen time Anxhelën, bashkëpunimi me të cilën ka përfaqësuar një mbështetje të rëndësishme profesionale dhe akademike. Diskutimet shkencore, shkëmbimi i ideve dhe puna e përbashkët kanë kontribuar në mënyrë të vlefshme gjatë këtij rrugëtimi kërkimor dhe akademik.

Në fund, dua të them se secili prej jush përfaqëson një pjesë të rëndësishme të jetës sime. Ashtu si pjesët e një puzzle-i që bashkohen për të krijuar një imazh të plotë e të bukur, edhe ju keni kontribuar që unë të ndërtoj jo vetëm këtë disertacion, por edhe një histori jete të mbushur me kuptim, dashuri dhe forcë.

Përjetësisht falënderuese ndaj jush!

Ilma

Parathënie

Zhvillimi i shpejtë i Inteligjencës Artificiale dhe transformimi dixhital i proceseve vendimmarrëse kanë ndryshuar mënyrën se si sistemet moderne analizojnë të dhënat, ndërveprojnë me përdoruesit dhe realizojnë procese automatike në mjedise gjithnjë e më komplekse. Kjo ka bërë që vendimmarrja automatike të mos konsiderohet më thjesht një koncept teorik, por një realitet teknologjik me ndikim të drejtpërdrejtë në fusha të ndryshme të jetës dhe industrisë. Pikërisht ky zhvillim dhe sfidat që e shoqërojnë kanë qenë motivimi kryesor për realizimin e këtij disertacioni.

Ky punim përfaqëson një përpjekje për të analizuar dhe kuptuar mënyrën se si teknikat e Inteligjencës Artificiale mund të kontribuojnë në realizimin e sistemeve të vendimmarrjes automatike, duke kombinuar performancën algoritmike me transparencën, interpretueshmërinë dhe besueshmërinë e sistemeve inteligjente. Në këtë kuadër, disertacioni trajton jo vetëm aspektet teorike të vendimmarrjes automatike dhe Inteligjencës Artificiale, por edhe aplikimin e tyre në kontekste konkrete si hoteleria, analiza e sentimentit nga e folura në gjuhën shqipe dhe monitorimi inteligjent i cilësisë së ajrit në ambiente të brendshme përmes konceptit të Binjakut Dixhital (Digital Twin).

Rrugëtimi për realizimin e këtij studimi ka qenë një proces i gjatë kërkimor dhe profesional, i shoqëruar me sfida të shumta që lidhen me analizën e të dhënave, ndërtimin e modeleve inteligjente dhe integrimin e sistemeve autonome në skenarë realë. Megjithatë, ky proces ka përfaqësuar një mundësi të vlefshme për të thelluar njohuritë në fushën e Inteligjencës Artificiale dhe për të eksploruar potencialin e saj në ndërtimin e sistemeve adaptive dhe të shpjegueshme.

Në një kohë kur sistemet inteligjente po bëhen gjithnjë e më të pranishme në vendimmarrjen organizative dhe shoqërore, ky disertacion synon të japë një kontribut modest në kuptimin dhe zhvillimin e sistemeve të automatizuara që jo vetëm performojnë me saktësi, por edhe janë të besueshme, transparente dhe të pranueshme për përdoruesit. Në këtë mënyrë, puna synon të ndërtojë një urë lidhëse ndërmjet zhvillimeve teorike të Inteligjencës Artificiale dhe aplikimeve praktike të saj në mjedise reale dhe dinamike.

Pasqyra përmbledhëse

Falënderime.....	V
<i>Parathënie</i>	VI
Lista e tabelave	V
Lista e figurave	VI
Lista e shkurtimeve	VIII
Hyrje	IX
KAPITULLI I Vështrim i përgjithshëm	1
1.1. Konteksti dhe rëndësia e vendimmarrjes automatike.....	1
1.2. IA në transformimin e proceseve vendimmarrëse.	2
1.3. Qëllimi dhe objektivat e studimit.....	4
1.4. Pyetjet kërkimore dhe hipotezat.....	5
1.5. Kontributi shkencor i pritur.....	7
KAPITULLI II Kuadri teorik dhe literatura ekzistuese	9
2.1. Bazat konceptuale të vendimmarrjes automatike	9
2.1.1. Llojet e Vendimmarrjeve Automatike	9
2.1.2. Evolucioni historik.....	10
2.1.3. Spektri Autonomisë dhe nivelet e ndërhyrjes njerëzore	11
2.1.4. Kuadri Rregullator si Pikë Referimi Konceptuale	11
2.1.5. Vendimmarrja automatike si sistem socio-teknik.....	12
2.2. Modelet teorike të vendimmarrjes.....	12
2.2.1. Paradigmat normative, deskriptive dhe preskriptive	13
2.2.2. Teoria e Racionaliteti të Kufizuar	13
2.2.3. Teoria e Perspektivës, heuristikat dhe bias-et.....	13
2.2.4. Teoria e Utilitetit të Pritur.....	14
2.2.5. Sintetizimi modeleve teorike dhe roli tyre në IA.....	14
2.3. Nga teoria në algoritëm: teknikat e IA në vendimmarrje automatike	15
2.3.1. Modelet probabilitike dhe vendimmarrja nën pasiguri	15
2.3.2. ML (Machine Learning) i orientuar drejt vendimeve	16
2.3.3. DL (Deep Learning) dhe vendimmarrja e bazuar në të dhëna të pastrukturuara	18

2.3.4. LLM (Large Language Model) dhe vendimmarrja gjeneruese.....	20
2.3.5. Inteligjenca Artificiale e Shpjegueshme si shtresë e detyrueshme	22
2.3.6. Sinteza: teknikat e IA sipas funksionit në ciklin e Vendimarrjes Automatike	23
2.3.7. Qasje të tjera të Inteligjencës Artificiale për vendimmarrje	24
2.4. Algoritmet Optimizues.....	24
KAPITULLI III Analiza e strukturës së të dhënave	26
3.1. Qëllimi i analizës së strukturës së të dhënave	26
3.2. Tipologjia e të dhënave dhe ndikimi në zgjedhjen e teknikave	26
3.3. Analiza e të dhënave të strukturuar nga sensorët mjedisor	27
3.3.1. Rregullsia kohore dhe cilësia e matjeve.....	27
3.3.2. Statistika përshkuese dhe shpërndarja e variablave.....	29
3.3.3. Analiza e pragjeve të lagështirës.....	30
3.3.4. Varësia kohore dhe veçoritë me vonesë kohore	31
3.3.5. Stacionariteti dhe sezonaliteti i serisë kohore	32
3.3.6. Implikime për modelim parashikues dhe vendimarrje automatike.....	34
3.4. Analiza e të dhënave të pa strukturuar audio	35
3.4.1. Struktura dhe shpërndarja e të dhënave audio	36
3.4.2. Ndashmëria e veçorive.....	39
3.5. Lidhja me vendimmarrjen automatike	45
KAPITULLI IV Integrimi i Shpjegueshmërisë së Inteligjencës Artificiale në ndihmë të vendimarrjes në hoteleri.....	46
4.1. Chatbot-et dhe automatizimet inteligjente në turizëm	47
4.1.1. Chatbot-et në mikëpritje	48
4.1.2. Chatbot-et gjeneruese dhe modelet e mëdha gjuhësore (LLMs)	50
4.1.3. IA e Shpjegueshme (XAI) dhe Sistemet e bashkëbisedimit	51
4.2. Chatbot Hibrid me XAI në Hoteleri.....	52
4.2.1. Pasqyra e përgjithshme e sistemit	52
4.2.2. Ndarja Funksionale në Shtresa e Sistemit.....	53
4.2.3. Moduli shpjegueshmërisë (Integrimi i XAI).....	54
4.2.4. Implementimi praktik dhe ndërfaqja.....	55
4.3. Metodologjia	57

4.4. Analiza e vlerësimit të modelit dhe të dhënave	59
4.5. Teknikat e balancimit.....	61
4.6. Analiza konkluzioneve.....	63
KAPITULLI V Analiza e Sentimentit nga e Folura në Gjuhën Shqipe: Një Qasje e të	
Mësuarit të Thellë për Skenarë me Burime të Kufizuara dhe vendimarrja automatike	
64	
5.1. Gjuha shqipe si një gjuhë me burime të kufizuara	64
5.2. Kategorizimi i karakteristikave akustike për njohjen e emocioneve	65
5.3. Metodologjia e nxjerrjes së karakteristikave për njohjen e emocioneve nga audio	66
5.4. Përgatitja datasetit	67
5.4.1. Strategjia e etiketimit të emocioneve të audios.....	70
5.4.2. Paraprocesimi.....	70
5.5. Ekstradimi i karakteristikave me MFCC dhe HuBERT.....	71
5.6. Klasifikues neural i lehtë për analizën e emocionit në audio.....	71
5.6.1. Rrjeti Feed Forward me dy shtresa	72
5.6.2. Modeli klasifikues Dy shtresor Feed Forward për MFCC dhe HuBERT... ..	72
5.6.3. LSTM dydrejtimore (Bi-LSTM) me Shtresën e vëmendjes	74
KAPITULLI VI Framework sugjerues për Vendimarrjen Automatike bazuar në	
Binjakun Dixhital.....	
80	
6.1. Prezantimi i problemit.....	80
6.2. Koncepti i Binjakut Dixhital	80
6.3. Sinkronizimi me sensorët real	81
6.4. Shtresa e përpunimit dhe ruajtjes së të dhënave.....	85
6.5. Integrimi modeleve të Inteligjencës Artificiale.....	88
6.5.1. Gjetja e intervalit kohor parashikues	88
6.5.2. Përcaktimi i pragut kritik të lagështirës relative	92
6.5.3. Aplikimi modeleve statistikore, të mësuarit të makinës dhe të mësuarit të thelluar	93
6.6. Autonomia e modulit të vendimarrjes.....	98
Përfundime.....	105
Punime për të ardhmen	107
Aneks -Kodi implementuar.....	108

Bibliografia	120
Lista e publikimeve	130

Lista e tabelave

Tabela 2.1 Kategorizimi i vendimarrjes automatike sipas literaturës bashkëkohore.....	10
Tabela 2.2 Teknikat e IA sipas funksionit në ciklin vendimarrjes automatike dhe nivelit të shpjegueshmërise.....	23
Tabela 3.1 Grupimi sipas muajve i datave dhe sasisë së matjeve të munguara.....	27
Tabela 3.2 Vlerësimi i anomalive sipas attributeve të dataset-it	28
Tabela 3.3 Statistika përshkruese e lagështirës.....	30
Tabela 3.4 Vlerësimi i stacionaritetit me testin ADF	33
Tabela 3.5 Statistika përshkruese të audios	35
Tabela 3.6 Analiza përmbledhëse e karakteristikave akustike të dataset-it audio para normalizimit	36
Tabela 3.7 Analiza e karakteristikave akustike të dataset-it audio pas aplikimit të normalizimit.....	36
Tabela 3.8 Statistika përmbledhëse e veçorive akustike për klasat emocionale.....	40
Tabela 3.9 Performanca e modelit të klasifikimit për secilën klasë	44
Tabela 5.1 Statistikë vlerësuese e dataset-it audio.....	67
Tabela 5.2 Distanca mesatare ndërmjet segmenteve	70
Tabela 5.3 Raporti Klasifikues	73
Tabela 5.4 Saktësia dhe Madhësia F1	74
Tabela 5.5 Raporti klasifikues.....	79
Tabela 6.1 Analiza e variacionit të tre horizonteve parashikues	89
Tabela 6.2 Rezultatet statistikore të horizonteve parashikuese	89
Tabela 6.3 Shpërndarja e lagështirës sipas muajve të vitit 2025	92
Tabela 6.4 Vlerësimi i modeleve për pragun fiks 69%	95
Tabela 6.5 Vlerësimi i modeleve për pragun interval 58-72%.....	96
Tabela 6.6 Krahësimi i performancës së modeleve të ML në dataset-in e trajnimit dhe testimit për kontrollin e mbipërshtatjes (overfitting).....	98
Tabela 6.7 Përshkrimi i rasteve të vendimarrjes.....	99
Tabela 6.8 Lidhja logjike e gjendjes aktuale me atë vendimarrjes	100
Tabela 6.9 Krahësimi i rasteve TP, TN, FP dhe FN për modelet parashikuese	102

Lista e figurave

Figura 2.1 Nivelet e automatizimit të vendimmarrjes dhe roli i algoritmit në ndërhyrje	11
Figura 2.2 Kurba e teorisë së perspektivës me vlerën subjektive të fitimeve dhe humbjeve	13
Figura 2.3 Vendimarrja sipas SEU me alternativa për maksimizimin e Utilitetit të Pritur	14
Figura 3.1 Heatmap autokorrelacionit për serinë kohore të lagështirës, temperaturës dhe trysnisë	29
Figura 3.2 Grafiku i shpërndarjes për tre parametrat	29
Figura 3.3 Shpërndarja e vlerave të lagështirës sipas frekuencës.....	30
Figura 3.4 Autokorrelacioni i lagështirës relative sipas vonesave kohore në intervale 5-minutëshe.....	32
Figura 3.5 Sezonaliteti ditor i lagështirës sipas orës	33
Figura 3.6 Sezonaliteti mujor i lagështirës	34
Figura 3.7 Komponenti trend i lagështirës relative gjatë vitit	34
Figura 3.8 Histogram i kohëzgjatjes.....	37
Figura 3.9 Histograma e pauzave në audio për tre klasat.....	38
Figura 3.10 RMS mesatare ndërmjet klasave pozitive, negative dhe neutrale	38
Figura 3.11 Shpërndarja e veçorive audio për vlerësimin e ndashmërisë së klasave	41
Figura 3.12 Dhjetë veçoritë me ndikim më të lartë	43
Figura 3.13 Aplikimi metodës PCA mbi të dhënat audio	43
Figura 3.14 Projektioni i të dhënave në 2D me t-SNE.....	45
Figura 4.1 Kuantifikimi i impaktit të IA dhe Automatizimit në fushën e mikpritjes	47
Figura 4.2 Cikli i jetës së chatbot-it	49
Figura 4.3 Përmbledhje e Sistemit Hibrid të Chatbot-it.....	53
Figura 4.4 Kategorizimi i inteligjencës artificiale të shpjegueshme	54
Figura 4.5 Ndërfaqja e përdoruesit.....	56
Figura 4.6 Përgjigje e gjeneruar nga Inteligenca Artificiale	56
Figura 4.7 Funkzioni për vlerësimin e përdoruesit.....	57
Figura 4.8 Paraqitja skematike e funksionit vendimarrës.....	58
Figura 4.9 Grupimet e funksionit për gjenerimet bazë.....	58
Figura 4.10 Funkzioni rishikimit të procesit.....	59
Figura 4.11 Raporti klasifikimit të dy klasave	60
Figura 4.12 Rezultatet e klasifikimit pas aplikimit të teknikave balancuese.....	61
Figura 5.1 Vlera mesatare për S.Flatness dhe ZCR Pozitive dhe Negative	68
Figura 5.2 Vlera mesatare për S.Flatness dhe ZCR Neutrale.....	68
Figura 5.3 Mesatarja e ZCR për klasat e audios.....	69
Figura 5.4 Funkzioni aktivizimit ReLU	72
Figura 5.5 Funkzioni shtresës Dense	72
Figura 5.6 Matrica e konfuzionit për MFCC për klasifikuesin MLP	73
Figura 5.7 Matrica e konfuzionit për HuBERT për klasifikuesin MLP	73
Figura 5.8 Perceptimi vizual i matricës me frame kohore dhe vektor.....	75
Figura 5.9 Kurba e humbjes gjatë trajnimit dhe validimit – BiLSTM + Attention + MFCC.....	77
Figura 5.10 Matrica konfuzionit për Bi-LSTM me Shtresën e Vëmendjes përpara class weight loss.....	78
Figura 5.11 Matrica e konfuzionit për Bi-LSTM me Shtresën e Vëmendjes pas weight class loss	78
Figura 5.12 Matrica e konfuzionit për HuBERT+BiLSTM me shtresën e vëmendjes	79
Figura 6.1 Framework-u i Binjakut Dixhital (DT) në ndihmesë të realizimit të vendimmarrjes.....	81
Figura 6.2 Pajisja AirVisual dhe rezultatet e aplikacionit online	82
Figura 6.3 Arkitektura e pajisjes Arduino dhe modelimi fizik.....	83

Figura 6.4 Struktura e veprimtarisë dhe pyetësi online.....	84
Figura 6.5 Platforma në cloud e IQAir.....	85
Figura 6.6 Krijimi i skedulimit për ekzekutimin e skriptit në Task Scheduler të Windows	86
Figura 6.7 Trigeri i detyrës automatik i konfiguruar çdo 5 min	86
Figura 6.8 Databaza SQLite	87
Figura 6.9 Skedulimi taskut ne SQLite	87
Figura 6.10 Metrika MAE për horizonte nga tre modelet	91
Figura 6.11 Metrika RMSE për horizonte nga tre modelet parashikuese	91
Figura 6.12 Metrika R^2 për horizonte nga tre modelet parashikuese	91
Figura 6.13 Krahasimi ndërmjet vlerave reale dhe parashikimeve të modelit Linear Regression për lagështirën relative.	97
Figura 6.14 Krahasimi ndërmjet vlerave reale dhe parashikimeve të modelit Pyllit te rastesishëm për lagështirën relative.	97
Figura 6.15 Krahasimi ndërmjet vlerave reale dhe parashikimeve të modelit XGBoost për lagështirën relative.	98
Figura 6.16 SQLite tabela e vendimeve	100
Figura 6.17 Krahasimi i Recall dhe F1-score për modelet parashikuese në rastin e pragut fiks dhe intervalit të vendimmarrjes.....	101
Figura 6.18 Krahasimi i Saktësisë dhe Precizionit për modelet parashikuese në rastin e pragut fiks dhe intervalit të vendimmarrjes.....	101
Figura 6.19 Thithësi lagështirës dhe priza inteligjente.....	103
Figura 6.20 Platforma e zhvilluesve Tuya.....	103

Lista e shkurtimeve

AI	Inteligjenca Artificiale
ACF	Funksioni Korrelacionit Automatik
AI Act	Ligj BE rregullon përdorimin e Inteligjencës Artificiale
ANN	Rrjetet Neurale Artificiale
BB	Kutia e zezë
BD	Sasia e madhe e të dhënave
CNN	Rrjet Nervor Konvolucional
DARPA	Agjencia e Projekteve të Avancuara të Kërkimit në Mbrojtje
DL	E mësuara e thelluar
DSS	Sistemet e Mbështetjes së Vendimmarrjes
DSS	Sistemet tradicionale të mbështetjes së vendimit
DT	Digital Twin
ES	Sistemet Eksperte
GAI	Inteligjenca Artificiale Gjeneruese
GDPR	Mbrojtja e të dhënave personale
GRU	Njësi Rekurente me Porta
HMM	Modelet e fshehta të Markovit
IAQ	Cilësia e Ajrit në Ambjente të Brendshme
IoT	Interneti i Gjërave
LIME	Shpjegime Lokale të Interpretueshme, të Pavarura nga Modeli
LLM	Modelet e Mëdha Gjuhësore
LR	Regresioni Linear
LSTM	Memorie afatshkurtër dhe afatgjatë
MAE	Gabimi Mesatar Absolut
MDP	Procesi Vendimmarrjes Markov
MFCC	Koeficientët cepstralë në shkallën Mel
ML	E mësuara e makinës
NDP	Procesi “nga ide në produkt final”
NLP	Përpunimi i gjuhës natyrore
PCA	Analiza Komponentëve Kryesor
RAG	Gjenerim i përforcuar nga kërkimi i informacionit
RF	Random Forest
RL	Të mësuarit përmes përforcimit
RMS	Rrënja mesatare katrore
RS	Sistemet e Rekomandimit
SER	Njohja e Emocioneve nga e Folura
SEU	Teoria e Utilitetit të Pritur
SHapley	Shpjegime Shtesore të bazuara në vlerat Shapley
SVM	Makina e Vektorëve Mbështetës
XAI	Inteligjenca Artificiale e Shpjegueshme
XNN	Rrjetet Neurale të Shpjegueshme

Hyrje

Zhvillimi i Inteligjencës Artificiale ka kaluar nëpër disa etapa të rëndësishme historike, duke evoluar nga konceptet teorike të inteligjencës së makinave deri te sistemet moderne të të mësuarit të thellë dhe vendimmarrjes automatike. Një nga punimet themelore në këtë fushë është artikulli i Alan Turing “Computing Machinery and Intelligence” (1950), ku u ngrit pyetja nëse makinat mund të mendojnë dhe u prezantua koncepti i “Turing Test”, i cili konsiderohet si një nga bazat teorike të Inteligjencës Artificiale. Turing argumentoi se sistemet kompjuterike mund të simulojnë sjellje inteligjente përmes ndërveprimit me njeriun, duke hapur rrugën për zhvillimet e mëvonshme në IA [1].

Më vonë, Inteligenca Artificiale kaloi nëpër disa faza zhvillimi, duke përfshirë sistemet simbolike, të mësuarit e makinës dhe të mësuarit e thelluar. Haenlein dhe Kaplan (2019) paraqesin një analizë historike të evolucionit të IA, duke theksuar periudhat e avancimit dhe të ngadalësimit të kësaj fushe, të njohura si “Dimrat e IA”. Autorët argumentojnë se rritja e fuqisë kompjuterike, disponueshmëria e të dhënave masive dhe zhvillimi i algoritmeve të avancuara kontribuan në rikthimin dhe përhapjen e IA moderne në fusha të ndryshme të industrisë dhe kërkimit shkencor [2].

Në dekadat e fundit, IA është transformuar në një paradigmë të rëndësishme për analizën inteligjente të të dhënave dhe automatizimin e proceseve vendimmarrëse. Në artikullin [3] theksohet se zhvillimi i teknikave të mësimit të thelluar, rrjeteve neurale dhe Të Dhënat e Mëdha ka rritur ndjeshëm aftësinë e sistemeve inteligjente për të analizuar të dhëna komplekse dhe për të mbështetur sisteme autonome dhe vendimmarrje automatike. Këto zhvillime kanë bërë të mundur aplikimin e IA në fusha të ndryshme, përfshirë analizën e sentimentit, sistemet inteligjente të rekomandimit, Binjaku Dixhital dhe menaxhimin inteligjent të mjediseve të brendshme.

Zhvillimi i IA ka mundur krijimin e sistemeve autonome të afta për të perceptuar mjedisin, për të analizuar të dhëna dhe për të marrë vendime pa ndërhyrje të vazhdueshme njerëzore. Automatizimi ka evoluar duke filluar nga sistemet tradicionale të bazuara në rregulla drejt sistemeve inteligjente adaptive që përdorin të mësuarin e makinës dhe të mësuarën e thelluar për të optimizuar proceset automatike në industri, robotikë, transport dhe sisteme inteligjente të vendimmarrjes [1].

Revolucioni i Inteligjencës Artificiale ka ndikuar në rritjen e sistemeve automatike dhe transformimin e proceseve organizative. Nga zhvillimi i algoritmeve inteligjente, fuqisë kompjuterike dhe disponueshmërisë së të dhënave është bërë i mundur automatizimi i proceseve që më parë kërkonin ndërhyrje njerëzore. Inteligenca Artificiale po kalon nga automatizimi i detyrave rutinë drejt sistemeve autonome që mund të analizojnë situata komplekse, të parashikojnë rezultate dhe të mbështesin vendimmarrjen në kohë reale [4].

Përparimet në Inteligjencën Artificiale, sistemet autonome dhe analiza e të dhënave kanë transformuar mënyrën se si operojnë organizatat dhe sistemet teknologjike. Automatizimi modern nuk kufizohet më vetëm në procese mekanike, por përfshin edhe analiza inteligjente, sisteme rekomandimi dhe platforma të automatizuara të vendimmarrjes që përshtaten në mënyrë dinamike ndaj mjedisit dhe të dhënave [5].

Inteligenca Artificiale përdoret kryesisht në tre drejtime kryesore: automatizimi i proceseve, analizimi inteligjent i të dhënave dhe ndërveprimi me përdoruesit. Sistemet IA ndihmojnë organizatat të analizojnë sasi të mëdha të dhënash, të identifikojnë modele dhe

të mbështesin vendimarrjen më të shpejtë e më të saktë. Kalimi nga sistemet tradicionale të vendimarrjes drejt sistemeve inteligjente autonome është bërë në mënyrë graduale. Inteligjenca Artificiale ka ndikuar edhe në strukturat moderne të vendimarrjes organizative. Ajo po transformon mënyrën se si merren vendimet, duke kaluar nga modele plotësisht njerëzore në sisteme hibride ku njeriu dhe algoritmet bashkëpunojnë në procese vendimarrës. Sistemet inteligjente janë të aftë të analizojnë të dhëna komplekse, të gjenerojnë rekomandime dhe në disa raste të realizojnë vendimarrje automatike mgjs sfida kryesore qëndron transparenca, besueshmëria dhe interpretueshmëria e IA në vendimarrje [6].

Pavarësisht çdo gjëje IA nuk zëvendëson plotësisht njeriun, por funksionon si një mekanizëm mbështetës që përmirëson cilësinë dhe shpejtësinë e vendimeve përmes analizës së të dhënave dhe parashikimeve inteligjente [7].

Rritja masive e të dhënave ka transformuar mënyrën se si organizatat analizojnë informacionin dhe marrin vendime. Të dhënat janë elementi kyç, për zhvillimin e sistemeve inteligjente dhe automatizimin e proceseve, duke mundësuar identifikimin e modeleve, parashikimin e sjelljeve dhe optimizimin e vendimarrjes në kohë reale [8]. Sistemet moderne të IA varen nga cilësia, volumi dhe përpunimi i të dhënave për të gjeneruar modele parashikuese dhe për të mbështetur procese adaptive automatike [9]. Automatizimi modern nuk kufizohet vetëm në procese mekanike, por përfshin analiza inteligjente, sisteme rekomadimi dhe platforma adaptive të bazuara në të dhëna [10].

Megjithëse automatizimi mund të rrisë efikasitetin dhe shpejtësinë e proceseve, ai krijon edhe sfida të rëndësishme si humbja e kontrollit njerëzor, mungesa e transparencës, besueshmëria e sistemeve dhe vështirësia për të kuptuar mënyrën se si sistemet automatike arrijnë në një vendim. Nivelet e larta të automatizimit mund të shkaktojnë varësi të tepërt nga sistemi dhe ulje të ndërgjegjësimit situacional të përdoruesit [11].

Sistemet moderne të të mësuarit të thelluar mund të arrijnë performancë shumë të lartë, por shpesh funksionojnë si “kuti e zezë”, ku përdoruesit nuk arrijnë të kuptojnë logjikën e brendshme të vendimit. Mungesa e transparencës krijon sfida në besueshmëri, përgjegjësi dhe përdorim të IA në fusha kritike si shëndetësia, financa dhe sistemet autonome të vendimarrjes [12].

Sistemet e Inteligjencës Artificiale duhet të integrojnë të dhëna nga burime të ndryshme, të cilat shpesh kanë formate, struktura dhe cilësi të ndryshme. Heterogjeniteti i të dhënave krijon probleme në integrim, pastrim, standardizim dhe analizë, duke ndikuar drejtpërdrejt në performancën e modeleve inteligjente dhe sistemeve automatike të vendimarrjes. Gjithashtu, evidentohet rëndësia e proceseve të “kujdesi ndaj të dhënave” dhe harmonizimit të të dhënave për ndërtimin e sistemeve të besueshme të IA [13].

Sfidat të mëdha hasen në përpunimin e gjuhëve me burime të kufizuara. Modelet moderne të NLP (Natural Language Processing) dhe të mësuarit të thelluar performojnë shumë mirë për gjuhët me burime të kufizuara ndërsa me më pak të dhëna përballen me probleme në klasifikim, analizë sentimentit dhe automatizim të proceseve gjuhësore. Kjo lidhet drejtpërdrejt me sfidat e analizës së sentimentit në gjuhën shqipe [14].

Për të përmirësuar performancën në gjuhët me burime të kufizuara propozohet përdorimi i teknikave të “transferimit të të mësuarës” dhe modeleve shumëgjuhëshe [15]. Ky koncept

lidhet edhe me zhvillimin e modeleve për analizën audio dhe sentimentin në gjuhën shqipe.

Rëndësia shkencore dhe praktike e Inteligjencës Artificiale lidhet me aftësinë e saj për të mbështetur dhe automatizuar proceset e vendimmarrjes në fusha të ndryshme aplikative. Sipas [16], sistemet moderne të IA përdoren jo vetëm për automatizimin e proceseve rutinë, por edhe për analizën inteligjente të të dhënave dhe gjenerimin e rekomandimeve vendimmarrëse në kohë reale. Zhvillimi i këtyre sistemeve ka ndikuar ndjeshëm edhe në sektorin e hotelarisë dhe turizmit, ku Inteligjenca Artificiale përdoret për personalizimin e shërbimeve, analizën e sjelljes së klientëve dhe përmirësimin e proceseve operacionale përmes sistemeve inteligjente dhe automatizimit të shërbimeve [17]. Paralelisht, përdorimi i IA në analizën e sentimentit dhe përpunimin e gjuhëve me burime të kufizuara ka krijuar mundësi të reja për interpretimin automatik të emocioneve dhe analizën e komunikimit njerëzor. Megjithatë, siç theksohet nga [14], gjuhët me më pak burime, si gjuha shqipe, vazhdojnë të përballen me sfida që lidhen me mungesën e dataset-eve dhe modeleve të trajnuara. Në të njëjtën kohë, zhvillimi i konceptit të Binjakut Dixhital dhe integrimi i Inteligjencës Artificiale me sensorët IoT kanë krijuar sisteme adaptive të afta për monitorim në kohë reale, analizë parashikuese dhe vendimmarrje automatike në menaxhimin e cilësisë së ajrit në ambiente të brendshme [18]. Këto zhvillime tregojnë se Inteligjenca Artificiale dhe vendimmarrja automatike po shndërrohen në komponentë kyç të sistemeve inteligjente moderne dhe aplikimeve adaptive të bazuara në të dhëna.

KAPITULLI I

Vështrim i përgjithshëm

1.1. Konteksti dhe rëndësia e vendimmarrjes automatike

Në dekadat e fundit, zhvillimi i shpejtë i teknologjive digjitale dhe rritja eksponenciale e sasisë së të dhënave kanë transformuar mënyrën se si organizatat, institucionet dhe individët marrin vendime. Proceset tradicionale të vendimmarrjes, të bazuara kryesisht në përvojën njerëzore dhe analiza manuale, po përballen me kufizime serioze përballë kompleksitetit, dinamikës dhe shkallës së sistemeve moderne. Në këtë kontekst, vendimmarrja automatike ka marrë një rol gjithnjë e më të rëndësishëm, duke u shndërruar në një komponent thelbësor të sistemeve inteligjente bashkëkohore.

Vendimmarrja automatike përkufizohet si procesi përmes të cilit një sistem kompjuterik, me ndërhyrje minimale ose pa ndërhyrje të drejtpërdrejtë njerëzore, analizon të dhënat, identifikon modele dhe prodhon vendime ose rekomandime konkrete [19], [20]. Kjo qasje mbështetet në teknika të inteligjencës artificiale dhe të të mësuarit makinerik, të cilat mundësojnë përpunimin e sasive të mëdha të të dhënave dhe reagimin në kohë reale ndaj ndryshimeve të mjedisit operacional.

Një nga faktorët kryesorë që ka nxitur zhvillimin e vendimmarrjes automatike është fenomeni i BD (Big Data). Karakteristikat e të dhënave moderne, vëllimi i madh, shpejtësia e gjenerimit dhe larmia e burimeve, e bëjnë të pamundur analizën efektive vetëm nga ekspertët njerëzorë [21]. Në këto kushte, sistemet e vendimmarrjes automatike ofrojnë mundësinë për të analizuar të dhënat në kohë reale, për të identifikuar modele të fshehura dhe për të ndërmarrë veprime të bazuara në evidencë. Në shumë fusha, si shëndetësia, industria, arsimi, hoteleria, transporti dhe menaxhimi i mjediseve të brendshme, vendimet duhet të merren në kohë reale dhe të bazohen në informacione që ndryshojnë vazhdimisht. Për këtë arsye, përdorimi i sistemeve të automatizuara për vendimmarrje nuk është më një zgjedhje opsionale, por një domosdoshmëri teknologjike dhe organizative.

Rëndësia e vendimmarrjes automatike lidhet drejtpërdrejt me aftësinë e saj për të rritur efikasitetin, saktësinë dhe shpejtësinë e reagimit. Modelet e inteligjencës artificiale, përfshirë metodat e ML (Machine Learning) dhe të DL (Deep Learning), mund të përdoren për të parashikuar sjellje, për të zbuluar anomali, për të klasifikuar gjendje dhe për të optimizuar vendimet në mënyrë adaptive. Ndryshe nga sistemet tradicionale të bazuara në rregulla fikse, këto modele janë në gjendje të mësojnë nga të dhënat historike dhe të përmirësojnë performancën e tyre me kalimin e kohës, duke e bërë vendimmarrjen automatike më fleksibile dhe më të përshtatshme për sisteme komplekse dhe dinamike. Megjithatë, rritja e përdorimit të vendimmarrjes automatike ka sjellë edhe sfida të rëndësishme shkencore, teknike dhe shoqërore [22], [23]. Shumë modele moderne, veçanërisht ato të bazuara në të mësuarit e thelluar, funksionojnë si sisteme “BB” (Black Box), ku procesi i brendshëm i marrjes së vendimit është i vështirë për t’u interpretuar. Kjo mungesë transparence krijon probleme në kontekstet ku vendimet kanë ndikim të drejtpërdrejtë në jetën e individëve. Për më tepër, modelet e inteligjencës artificiale mund të trashëgojnë paragjykime nga të dhënat e trajnimit, duke prodhuar vendime të padrejta ose diskriminuese. Këto sfida theksojnë nevojën për zhvillimin e sistemeve që jo vetëm janë të sakta, por edhe të drejtë, të besueshme dhe të kuptueshme.

Në këtë kuadër, koncepti i inteligjencës artificiale të shpjegueshme XAI (Explainable Artificial Intelligence) merr një rëndësi të veçantë. Integrimi i mekanizmave të shpjegueshmërisë, së bashku me mbikëqyrjen njerëzore dhe parimet etike, është thelbësor për të garantuar transparencën dhe besueshmërinë e sistemeve të vendimmarrjes automatike. Në nivel rregullator, veçanërisht në Bashkimin Evropian, theksohet e drejta e individëve për të kuptuar logjikën e vendimeve automatike që i prekin ata, duke rritur kërkesat për sisteme të përgjegjshme dhe transparente [24], [25].

Në këtë kontekst, vendimmarrja automatike duhet të trajtohet si një fushë ndërdisiplinore që lidh inteligjencën artificiale, shkencën e të dhënave, etikën dhe kornizat ligjore. Rëndësia e saj nuk qëndron vetëm në përmirësimin e efikasitetit operacional, por edhe në ndërtimin e sistemeve inteligjente që janë të pranueshme dhe të besueshme në praktikë.

Në këtë disertacion, vendimmarrja automatike trajtohet si një qasje ndërdisiplinore që lidh inteligjencën artificiale me aplikime konkrete në kontekste të ndryshme, si hoteleria, analiza e sentimentit nga e folura dhe menaxhimi i cilësisë së ajrit të brendshëm përmes konceptit të Digital Twin. Këto raste studimore demonstrojnë se vendimmarrja automatike mund të realizohet në forma të ndryshme, duke përfshirë rekomandime të shpjegueshme për përdoruesit, interpretimin e emocioneve nga sinjalet audio dhe aktivizimin e veprimeve automatike bazuar në parashikime mjedisore. Në këtë mënyrë, vendimmarrja automatike shndërrohet në një mekanizëm kyç për transformimin e të dhënave në vendime praktike, të matshme dhe të dobishme.

1.2. IA në transformimin e proceseve vendimmarrëse.

Termi Inteligjencë artificiale dhe sisteme të inteligjencës artificiale u prezantuan për herë të parë në vitet 1950. Që atëherë, IA ka përjetuar uljet e ngritjet e saja [26]. Proceset e vendimmarrjes janë bërë gjithnjë e më komplekse në sektorë si biznesi, kujdesi shëndetësor, financa dhe teknologjia, gjë që çoi në zhvillimin e mjeteve të sofistikuar për të rritur efikasitetin e vendimmarrjes. DSS (Decision Support System) të zhvilluara në vitet 1960–70, integruan automatizimin dhe përpunimin e të dhënave për të rritur vendimmarrjen njerëzore. Pas tyre, ES (Expert System) u zhvilluan në vitet 1970, duke nxjerrë njohuri nga ekspertët dhe duke i ruajtur si fakte dhe heuristika në një bazë njohurish, në mënyrë të ngjashme me një ekspert njerëzor. Më pas, RS (Recommendation system) shndërruan ndërveprimet dixhitale duke analizuar preferencat e përdoruesve për të ofruar rekomandime të personalizuar. Për shkak të disponueshmërisë së madhe të grupeve të të dhënave dhe fuqisë së lartë kompjuterike, sistemet e ML (Machine Learning) tani po performojnë mirë, por mungesa e interpretueshmërisë, besimit dhe transparencës çoi në emergjencën e XAI (eXplainable AI) si faza e fundit e evolucionit të vendimmarrjes [27].

Inteligjenca artificiale ka luajtur një rol transformues në evolucionin e proceseve vendimmarrëse, duke zhvendosur fokusin nga vendimmarrja tradicionale e bazuar në përvojë njerëzore drejt vendimmarrjes së bazuar në të dhëna dhe modele algoritmike. Studimet shkencore theksojnë se IA ka aftësinë të analizojë sasi të mëdha dhe komplekse të dhënash, të identifikojë modele jo-lineare dhe të prodhojë vendime ose rekomandime me një nivel saktësie dhe shpejtësie që tejkalon kapacitetet njerëzore në shumë kontekste operative (Russell & Norvig, 2021). Literatura shkencore tregon se Inteligjenca Artificiale

ka transformuar vendimmarrjen nga një proces reaktiv në një proces parashikues dhe proaktiv. Modelet parashikuese të bazuara në inteligjencë artificiale mundësojnë identifikimin e tendencave, vlerësimin e rreziqeve dhe simulimin e skenarëve të ndryshëm vendimmarrës përpara se vendimi të zbatohet në praktikë. Kjo qasje ka gjetur zbatim të gjerë në fusha si financa, shëndetësia, menaxhimi mjedisor dhe sistemet e rekomandimit (Shmueli & Koppius, 2011). Artikulli shkencor [28] paraqet se si inteligjenca artificiale ndikon në proceset e vendimmarrjes menaxheriale në menaxhimin e projekteve. Ajo sjell përmirësimin e cilësisë dhe konsistencës së të dhënave duke paraqitur vendimarrje më të shpejtë dhe efektive, por edhe duke identifikuar rreziqet si ndihmesë për menaxherët për të marrë vendime strategjike.

Proceset e vendimmarrjes mund të optimizohen me ndihmën e disa teknikave të Inteligjencës Artificiale të avancuara si ML (Machine Learning) apo rasti i përpunimit të gjuhës natyrore si dhe kur është e nevojshme analitika parashikuese, duke rritur ndjeshëm saktësinë, efikasitetin dhe përgjigjen e shpejtë në sektorë të ndryshëm si kujdesi shëndetësor, në fushën e financës, si dhe në administratën publike [29]. Pavarësisht këtyre përparimeve, identifikohen sfida që janë të vazhdueshme, si paragjykimi që ekziston mbi algoritmet, shqetësimet që lidhen me privatësinë e të dhënave dhe mungesën e transparencës në modelet BB (Black Box) të IA, dhe këto janë ato që dëmtojnë besimin.

Një fushë tjetër është edhe bujqësia ku një problem i rëndësishëm është cilësia e vendimeve menaxheriale. Në procesin e vendimmarrjes, Inteligjenca Artificiale konkretisht ndikon në diagnostifikimin e hershëm të problemeve, mbledhjen e analizimin e të dhënave, në reduktimin e pasigurisë, gjenerimin e alternativave optimale dhe në automatizimin dhe monitorimin [30]

Tabela 1.1 Analizë krahasuese e tre literaturave për procesin e vendimmarrjes në mjekësi

PËRSHKRIMI	[31] THE POTENTIAL FOR ARTIFICIAL INTELLIGENCE IN HEALTHCARE AI"	[32] HIGH-PERFORMANCE MEDICINE: THE CONVERGENCE OF HUMAN AND AI	[33] EXPLAINABLE AI: CONCEPTS, TAXONOMIES & CHALLENGES
FOKUSI KRYESOR	Potenciali i IA në sektorin e kujdesit shëndetësor: diagnostikim, trajtim dhe aktivitete administrative	Konvergjenca e inteligjencës njerëzore dhe IA për mjekësinë me performancë të lartë	Nevoja për Inteligjencë Artificiale të shpjgueshme (XAI) si parakusht për vendimarrje të besueshme dhe të përgjegjshme
SI TRANSFORMOHET PROCESI I VENDIMARRJES	Rekomandime diagnozë/trajtimi	Interpretim i shpejtë i imazheve mjekësore	Shpjegueshmëria e modeleve IA si kusht vendimarrjeje
	Angazhim i pacientit	Reduktim i gabimeve mjekësore	Kalon nga "kutia e zezë" në IA transparente
	Automatizim i aktiviteteve administrative	Mundëson pacientët të procesojnë të dhënat e tyre	Mbështet llogaridhënien dhe drejtësinë algoritmike
	IA kryen detyra po aq mirë ose më mirë se njerëzit në disa raste	Ndikim në 3 nivele: klinikistë, sisteme shëndetësore, pacientë	Mundëson zbatim të gjerë organizacional
TEKNOLOGJITË KYÇE	Machine learning, rrjete neurale artificiale, NLP, robotikë, IA fizike	Deep learning, big data i etiketuar, fuqi kompjuterike e lartë, ruajtja në cloud	Modele ML transparente, teknika post-hoc, rrjete neurale të shpjgueshme (XNN)
KONTRIBUTI UNIK PËR VENDIMARRJEN	IA nuk zëvendëson mjekun, por transformon kategorikisht çdo fazë të rrjedhës vendimmarrëse klinike	Vendimarrja nuk është vetëm mjekësore, pacienti bëhet aktor aktiv i procesit falë Inteligjencës Artificiale	Vendimarrja e mbështetur në Inteligjencë Artificiale kërkon transparencë, pa shpjgueshmëri, sistemi nuk mund të miratohet apo besohet
SFIDAT E TRANSFORMIMIT	Faktorë zbatimi pengojnë automatizimin e gjerë; rreziku i erozionit të marrëdhënies mjek-pacient	Paragjykimi algoritmik, privatësia dhe siguria e të dhënave, mungesa e transparencës	Balancimi i saktësisë me interpretueshmërinë; varësia nga të dhëna cilësore; sfida etike dhe ligjore

KONKLuzion PËR TRANSFORMIMIN	IA do të aplikohet gjithnjë e më shumë; zbatimi i kujdesshëm është çelësi i suksesit	Përmirësime të ndjeshme në saktësi, produktivitet dhe rrjedhë pune janë të realizueshme afatgjatë	Responsible A, Inteligjenca Artificiale e drejtë, e shpjgueshme dhe llogaridhënëse, është kusht i domosdoshëm për vendimmarrje
---	--	---	--

Në tabelën 1.1 artikulli i autorëve Davenport & Kalakota e shikon transformimin nga këndvështrimi i aplikimit praktik ku Inteligjenca Artificiale ndryshon çdo fazë të kujdesit, por zbatuesi njerëzor mbetet qendror. Topol shkon më thellë duke argumentuar se vendimmarrja bëhet shumëdimensionale jo vetëm mjeku, por edhe pacienti dhe sistemi janë transformuar. Arrieta et al. ndërkaq vendos kushtin paraprak të gjithë transformimit: nëse Inteligjenca Artificiale nuk është e shpjgueshme dhe transparente, vendimmarrja e bazuar në të nuk mund të miratohet apo besohet nga asnjë aktor.

Një tjetër aspekt thelbësor i rolit të Inteligjencës Artificiale është shkallëzimi i vendimmarrjes. Ndryshe nga vendimmarrja njerëzore, e cila është e kufizuar nga koha dhe kapaciteti kognitiv, sistemet e inteligjencës artificiale mund të marrin vendime në shkallë të gjerë dhe në kohë reale, duke përpunuar flukse të vazhdueshme të dhënash. Kjo veçori është veçanërisht e rëndësishme në kontekstin e BD (Big Data) dhe sistemeve të bazuara në sensorë dhe platforma digjitale (Provost & Fawcett, 2013).

Në përmbledhje, artikujt shkencorë tregojnë se inteligjenca artificiale ka transformuar rrënjësisht proceset vendimmarrëse duke i bërë ato më të shpejta, më të sakta dhe më të bazuara në të dhëna. Megjithatë, ky transformim kërkon një qasje të balancuar që kombinon fuqinë analitike të Inteligjencës Artificiale me mekanizma shpjgueshmërie dhe kontrolli, në mënyrë që vendimmarrja automatike të jetë jo vetëm efikase, por edhe e besueshme dhe e pranueshme në kontekstet reale të përdorimit.

1.3. Qëllimi dhe objektivat e studimit.

Në këtë kontekst, evidentohet nevoja për një studim të integruar që analizon jo vetëm performancën e teknikave të inteligjencës artificiale në vendimmarrje, por edhe kushtet e transparencës dhe besueshmërisë që e bëjnë këtë vendimmarrje të pranueshme në praktikë. Megjithëse literatura ekzistuese ka trajtuar veçmas aspektet teknike të algoritmeve dhe çështjet etike të sistemeve automatike, mbetet një boshllëk i dukshëm në integrimin e të dyjave brenda një qasjeje të vetme metodologjike, e zbatuar në kontekste të ndryshme dhe të verifikueshme empirikisht. Ky studim synon të adresojë pikërisht këtë boshllëk.. Në mënyrë më specifike, objektivat e këtij studimi janë si më poshtë:

- Të analizojë konceptet themelore të vendimmarrjes automatike dhe evolucionin e tyre në kontekstin e sistemeve të bazuara në inteligjencë artificiale.
- Të shqyrtojë dhe të krahasojë teknika të ndryshme të inteligjencës artificiale dhe ML (Machine Learning) të përdorura në proceset vendimmarrëse, duke vlerësuar performancën dhe kufizimet e tyre.
- Të vlerësojë ndikimin e modeleve të inteligjencës artificiale në cilësinë, shpejtësinë dhe qëndrueshmërinë e vendimeve automatike në mjedise komplekse dhe dinamike.
- Të analizojë sfidat që lidhen me mungesën e transparencës dhe interpretueshmërisë së modeleve vendimmarrëse automatike, veçanërisht në kontekstet ku vendimet kanë ndikim të drejtpërdrejtë mbi individët.
- Të eksplorojë rolin e metodave të XAI (Explainable Artificial Intelligence) në rritjen e shpjgueshmërisë, besimit dhe përgjegjshmërisë së sistemeve të vendimmarrjes

automatike.

- Të propozojë një qasje metodologjike ose model konceptual për integrimin e teknikave të inteligjencës artificiale dhe shpjegueshmërisë në ndërtimin e sistemeve vendimmarrëse automatike.

Në përmbledhje, ky studim synon të ndërtojë një urë midis zhvillimeve teknike të inteligjencës artificiale dhe kërkesave praktike, etike dhe rregullatore të vendimmarrjes automatike. Përmes analizës së teknikave ekzistuese dhe integritit të qasjeve shpjegueshmërisë, disertacioni synon të ofrojë një kontribut shkencor që mbështet zhvillimin e sistemeve inteligjente të besueshme, transparente dhe të qëndrueshme.

1.4. Pyetjet kërkimore dhe hipotezat.

Në literaturën bashkëkohore, vendimmarrja automatike e bazuar në inteligjencë artificiale është identifikuar si një nga faktorët kryesorë të transformimit digjital në organizata dhe sisteme komplekse [26], [34]. Megjithatë, studimet ekzistuese theksojnë se, krahas përfitimeve në performancë dhe efikasitet, mbeten sfida të rëndësishme që lidhen me transparencën, interpretueshmërinë dhe besueshmërinë e këtyre sistemeve [22], [35]. Në këtë kontekst, ky studim synon të adresojë boshllëqet ekzistuese duke analizuar rolin e teknikave të inteligjencës artificiale në vendimmarrjen automatike në mjedise të ndryshme aplikative.

Për të përcaktuar qartë drejtimin e studimit dhe për tu fokusuar në probleme konkrete që kërkojnë zgjidhje shkencore thelbësore janë pyetjet kërkimore. Ato ndihmojnë në lidhjen logjike të pjesëve kryesore të punimit duke filluar nga kuadri teorik, te metodologjia dhe deri te rezultatet. Me anë të tyre tregohet se çfarë synohet të arrihet dhe si do të vlerësohet kontributi kërkimit. Pyetjet e mëposhtme kanë marrë shkas nga literatura ekzistuese, duke u bazuar në boshllëqet, kufizimet dhe sfidat e identifikuara në studimet e mëparshme. Ato nuk janë krijuar në mënyrë të izoluar, por janë ndërtuar mbi njohuri aktuale shkencore për të sjellë një kontribut të ri dhe të justifikuar në fushën e studimit.

Pyetje 1

Cilat teknika të inteligjencës artificiale janë më efektive për ndërtimin e sistemeve të vendimmarrjes automatike në mjedise komplekse dhe dinamike?

Pyetje 2

Si ndikon përdorimi i inteligjencës artificiale në cilësinë, shpejtësinë dhe shkallëzimin e vendimeve krahasuar me qasjet tradicionale?

Pyetje 3

Në çfarë mase modelet e inteligjencës artificiale përballen me sfida të transparencës dhe interpretueshmërisë?

Pyetje 4

Si ndikon integrimi i metodave XAI (eXplainable Artificial Intelligence) në besimin dhe pranueshmërinë e sistemeve vendimmarrëse?

Pyetje 5

A mund të kombinohen performanca dhe shpjgueshmëria për të përmbushur kërkesat etike dhe rregullatore?

Ndërsa përsa i takon hipotezave bazuar në pyetjet kërkimore dhe në rishikimin e literaturës shkencore, ky studim propozon hipotezat e mëposhtme:

Hipoteza 1.

Modelet e bazuara në inteligjencë artificiale ofrojnë performancë më të lartë në vendimmarrje automatike krahasuar me qasjet tradicionale të bazuara në rregulla, në përputhje me gjetjet e studimeve të mëparshme në literaturë [26].

Hipoteza 2.

Përdorimi i inteligjencës artificiale përmirëson ndjeshëm shpejtësinë dhe shkallëzimin e vendimeve në mjedise me të dhëna dinamike, duke mundësuar analiza në kohë reale [34].

Hipoteza 3.

Modelet komplekse të bazuara në të mësuarën e thelluar karakterizohen nga një nivel më i ulët interpretueshmërie, duke rritur perceptimin e tyre si sisteme “kuti e zezë” [22] .

Hipoteza 4.

Integrimi i metodave të XAI (Explainable Artificial Intelligence) rrit transparencën dhe besimin e përdoruesve ndaj sistemeve vendimmarrëse automatike [35].

Hipoteza 5.

Sistemet që kombinojnë performancën e lartë me shpjgueshmërinë janë më të përshtatshme për aplikime në kontekste të ndjeshme dhe të rregulluara [24] .

Në funksion të ndërtimit të një qasjeje të integruar kërkimore, është zhvilluar një tabelë sintezë që paraqet në mënyrë të strukturuar marrëdhëniet ndërmjet pyetjeve kërkimore, hipotezave, kapitujve të disertacionit, rasteve studimore, metodologjive të përdorura dhe rezultateve të pritshme. Kjo tabelë shërben si një mekanizëm konceptual dhe operacional që siguron lidhjen logjike ndërmjet komponentëve teorikë dhe aplikativë të studimit, duke mundësuar një analizë të sistematizuar dhe një validim të qëndrueshëm të hipotezave në kontekste të ndryshme reale.

Tabela 1.2 Shpërndarja e pyetjeve kërkimore dhe hipotezave sipas kapitujve të disertacionit

<i>Pyetja kërkimore</i>	<i>Hipoteza</i>	<i>Kapitulli</i>	<i>Rasti studimor</i>	<i>Metodologjia</i>	<i>Rezultati</i>
<i>P1</i>	H1	II, III, IV, V, VI	Të tre rastet	Krahasim modelesh	Identifikim teknikash
<i>P2</i>	H1, H2	IV, V, VI	Chatbot, E folura, DT	Analizë performance	Rritje efikasiteti
<i>P3</i>	H3	II, IV	Chatbot (XAI)	Analizë interpretuese	Evidencë “black-box”
<i>P4</i>	H4	IV	Chatbot	SHAP/LIME	Rritje besimi
<i>P5</i>	H5	II, IV, VI	DT + Chatbot	Analizë etike + sistem real	Qasje e balancuar

Siç paraqitet në Tabelën 1.2, ekziston një lidhje e drejtpërdrejtë ndërmjet pyetjeve kërkimore, hipotezave dhe kapitujve përkatës të disertacionit. Secila hipotezë testohet përmes analizave teorike dhe implementimeve praktike në rastet studimore, duke siguruar një qasje të integruar për vlerësimin e rolit të inteligjencës artificiale në vendimmarrjen automatike. Kjo strukturë mundëson jo vetëm validimin e hipotezave, por edhe demonstrimin e aplikueshmërisë së tyre në kontekste reale.

1.5. Kontributi shkencor i pritur.

Nga ana konceptuale kontributi bazohet në zgjerimin e konceptit të vendimarrjes automatike duke e trajtuar atë si një proces shumë dimensional që kombinon analizën e të dhënave, modelimin parashikues dhe ekzekutimin e vendimeve në kohë reale. Ndryshe nga qasjet tradicionale që fokusohen vetëm në optimizim, ky punim demonstron se inteligjenca artificiale mund të përdoret për të ndërtuar sisteme autonome vendimarrëse të afta të operojnë në mjedise të ndryshme aplikative. Nga ana metodologjike janë paraqitur tre raste studimore të ndryshme. Rasti i parë është përdorimi i eXplainable AI (XAI) në chatbot-et inteligjente për të rritur transparencën dhe interpretueshmërinë e vendimeve të marra. Rasti i dytë është aplikimi i modeleve ML (Machine Learning) për analizën e emocioneve nga e folura në gjuhë me burime të kufizuara si gjuha shqipe. Rasti i tretë është integrimi i DT (Digital Twin) me modele parashikuese dhe modul vendimarrjeje për menaxhimin e cilësisë së ajrit në ambiente të brendshme. Kjo qasje tregon se teknikat e ndryshme të AI (Artificial Intelligence) mund të kombinohen në mënyrë efektive për të adresuar probleme komplekse vendimarrëse në kontekste heterogjene.

Një nga kontributet kryesore të këtij punimi është demonstrimi i faktit që modelet dhe teknikat e inteligjencës artificiale për vendimarrje automatike janë të transferueshme ndërmjet domajeve të ndryshme. Pra mekanizmat e vendimarrjes automatike, parimet parashikuese dhe vlerësuese mund të aplikohen në shërbime, në perceptimin njerëzor emocional si dhe në sisteme fizike. Nga ana teknike ofrohen implementime konkrete të sistemeve të bazuara në inteligjencë artificiale që realizojnë vendimarrje bazuar në kombinimin e rregullave dhe modeleve parashikuese, analizim i të dhënave audio për

klasifikimin emocionale si dhe integrim të sensorëve dhe modeleve në një arkitekturë DT (Digital Twin). Ana teknike është e zgjeruar me kodin dhe sistemin në Aneksin e këtij punimi. Gjatë studimit demonstron rëndësia e integritit të XAI (eXplainable Artificial Intelligence) në sisteme vendimarrëse duke rritur transparencën dhe besimin e përdoruesit, duke justifikuar vendimet, duke zbërthyer konceptin e BB (Black Box). Rëndësi të veçantë ka zhvillimi i metodave për analizën e emocioneve në gjuhën shqipe duke demonstruar se modelet DL (Deep Learning) mund të funksionojnë edhe në datasete të kufizuara, duke hapur mundësi për aplikime të reja në gjuhë jo të mbështetura gjerësisht. Përsa i takon sistemeve inteligjente propozohet një arkitekturë funksionale për integrimin e sensorëve IoT (Internet of Things), modeleve të inteligjencës artificiale, modulit ADM (Automated Decision Making) dhe sistemit të veprimit në një DT (Digital Twin) operativ. Kjo sjell rritje edhe të konfortit në menaxhimin e situatave kur dora e njeriut mundon si dhe rritjen e efikasitetit energjitik. Pra thelbi i këtij punimi bazuar në inteligjencë artificiale tregon se mund të:

- automatizohen vendime në mënyrë të besueshme
- realizohen transferime në domene të ndryshme
- funksionojë edhe në kushte jo ideale si LRL (low-resource language)
- kombinohen perceptimi (audio), analiza (të dhënat) dhe kontrolli i IoT (Internet of Things)
- integrohen në sisteme reale si DT (Digital Twin)
- bëhen më transparente me anë të XAI (eXplainable AI)

Ky studim demonstroi se inteligjenca artificiale nuk shërben vetëm për analizë apo optimizim, por përbën një mekanizëm të plotë për ndërtimin e sistemeve autonome vendimarrëse, të afta të operojnë në mënyrë të besueshme, të shpjegueshme dhe të adaptueshme në kontekste të ndryshme reale.

KAPITULLI II

Kuadri teorik dhe literatura ekzistuese

2.1. Bazat konceptuale të vendimarrjes automatike

Vendimarrja automatike, ADM (Automated Decision Making), i referohet përdorimit të algoritmeve, modeleve kompjuterike dhe sistemeve të bazuara në të dhëna me qëllim mbështetjen ose zëvendësimin e plotë ose të pjesshëm të vendimeve që merren nga njerëzit. Në literaturë, ADM trajtohet si proces ku një sistem përpunon të dhëna hyrëse, aplikon rregulla ose modele analitike dhe prodhon një rezultat që ndikon në zgjedhje, rekomandime ose veprime konkrete. Studimet e fundit e shohin ADM jo vetëm si teknologji, por si fenomen socio-teknik që ndikon në organizata, shërbime publike, ekonomi dhe jetën e përditshme [36].

Një dallim i rëndësishëm konceptual është ai ndërmjet automatizimit të bazuar në rregulla dhe inteligjencës artificiale të bazuar në të dhëna. Sistemet e bazuara në rregulla përdorin kushte të paracaktuara nga ekspertët, si “nëse ndodh X, atëherë kryej Y”. Këto sisteme janë të dobishme kur problemi është i qartë, rregullat janë të qëndrueshme dhe vendimi mund të shpjegohet lehtësisht. Megjithatë, kufizimi i tyre kryesor është mungesa e përshtatshmërisë ndaj situatave të reja. Në të kundërt, modelet e inteligjencës artificiale, veçanërisht ML (Machine Learning) dhe DL (Deep Learning), mësojnë modele nga të dhënat dhe mund të identifikojnë marrëdhënie komplekse që nuk janë të programuara drejtpërdrejt nga njeriu [37].

Historikisht, zhvillimi i vendimarrjes automatike lidhet ngushtë me evolucionin e inteligjencës artificiale. Në fazat e hershme dominuan sistemet eksperte, të cilat përpiqeshin të imitonin arsyetimin e ekspertëve njerëzorë përmes bazave të njohurive dhe rregullave logjike. Më pas, me rritjen e kapaciteteve kompjuterike dhe disponueshmërinë e të dhënave, fokusi u zhvendos drejt modeleve të mësimit makinerik, ku vendimi nuk bazohet vetëm në rregulla të shkruara manualisht, por në modele të nxjerra nga të dhënat. Në qasjet moderne, sistemet inteligjente trajtohen si agjentë racionalë që perceptojnë mjedisin, analizojnë gjendjen aktuale dhe zgjedhin veprime për të arritur objektiva të caktuara [38].

Në këtë kuptim, ADM mund të shihet si një vazhdimësi: nga sistemet e thjeshta të automatizimit me rregulla, drejt sistemeve prediktive të bazuara në të dhëna dhe më tej drejt agjentëve inteligjentë që marrin vendime në mënyrë dinamike. Për një disertacion që trajton teknikat e inteligjencës artificiale për vendimarrje automatike, ky dallim është thelbësor, sepse ndihmon të sqarohet se jo çdo automatizim është inteligjencë artificiale. Automatizimi me rregulla zbaton logjikë të paracaktuar, ndërsa IA synon të mësojë, të përshtatet dhe të përmirësojë vendimarrjen në bazë të të dhënave, kontekstit dhe objektivave të sistemit.

2.1.1. Llojet e Vendimarrjeve Automatike

Në literaturën bashkëkohore, vendimarrja automatike, nuk trajtohet si një koncept

homogjen, por si një spektër që varion nga sisteme mbështetëse të thjeshta deri te sisteme plotësisht autonome. Ky dallim është thelbësor për të kuptuar rolin real të inteligjencës artificiale në procesin e vendimmarrjes dhe nivelin e ndërhyrjes njerëzore në sisteme të ndryshme.

Një nga qasjet më të njohura për klasifikimin e niveleve të automatizimit është ajo e [39] të cilët propozojnë një shkallë me disa nivele automatizimi bazuar në ndërveprimin njeri–makinë. Kjo shkallë është përshtatur dhe thjeshtuar gjerësisht në literaturën moderne, duke u përmbledhur në tre kategori kryesore operative të ADM-së: sisteme sugjeruese [40], [41], sisteme mbështetëse [42], [43] dhe sisteme zëvendësuese [44].

Këto kategori përfaqësojnë një evolucion gradual nga vendimmarrja e drejtuar nga njeriu drejt vendimmarrjes së automatizuar, ku algoritmet kalojnë nga një rol këshillues në një rol vendimmarrës të pavarur. Ky tranzicion është veçanërisht i rëndësishëm në kontekstin e sistemeve moderne të bazuara në inteligjencë artificiale, ku automatizimi nuk kufizohet vetëm në analizën e të dhënave, por shtrihet edhe në kontrollin dhe veprimin në kohë reale.

Në vijim, paraqitet një klasifikim i tre llojeve kryesore të ADM-së (Tabela 2.1), duke theksuar rolin e algoritmit, nivelin e ndërhyrjes njerëzore dhe shembuj konkretë nga aplikime reale.

Tabela 2.1 Kategorizimi i vendimmarrjes automatike sipas literaturës bashkëkohore

Lloji	Roli i algoritmit	Shembull
Sugjerues (Nivel 1)	Rekomandon; vendimi final mbetet tek operatori njerëzor	Netflix, GPS, Spotify
Mbështetës (Nivel 2)	Zbarkues i ngarkesës njerëzore; njeriu miraton ose refuzon rezultatin algoritmik	Diagnoza mjekësore me IA, triazhi spitalor (IBM Watson, CAD radiologji)
Zëvendësues (Nivel 3)	Vendos plotësisht vetë, pa ndërhyrje njerëzore në kohë reale	Tregtimi algoritmik financiar (HFT), automjetet autonome (SAE L4)

2.1.2. Evolucion i historik

Kalimi nga sistemet e para të bazuara në rregulla, drejt algoritmeve të të mësuarit nga të dhënat dhe më tej drejt modeleve të avancuara të inteligjencës artificiale gjeneruese, janë kaluar disa faza duke sjellë përmirësime të rëndësishme në automatizimin dhe kompleksitetin e vendimmarrjes. Me anë të kësaj ndarjeje qartësohet edhe progresi nga sisteme statike dhe të kufizuara drejt sistemeve dinamike, adaptive dhe të afta për arsyetim të avancuar. Në vijim paraqiten katër faza kryesore të këtij evolucioni.

- Faza 1 (1970-1985) Sistemet eksperte dhe logjika e bazuar në rregulla
Sistemet eksperte përdornin rregulla “nëse-atëherë” për të simuluar arsyetimin njerëzor. Sisteme si MYCIN dhe DENDRAL paraqitën potencialin e Inteligjencës Artificiale por ishin të kufizuara në shkallëzim dhe fleksibilitet [45].
- Faza 2 (1985-2000) Logjika e ngurtë: pemë vendimesh dhe automatizimi robotik
Algoritmet ishin të strukturuar si pemët e vendimeve. Vendimmarrja mbetej e paracaktuar (“hard-coded”) dhe nuk përfshinte të mësuarën nga të dhënat .
- Faza 3 (2000-2015) ML (Machine Learning) dhe BD (Big Data)
Kjo ishte faza e kalimit të modelet që mësojnë nga të dhënat, duke mundësuar analiza

parashikuese dhe sisteme rekomandimi. Rritja e sasisë së të dhënave dhe fuqisë llogaritëse rriti ndjeshëm performancën dhe përdorimin praktik të Inteligjencës Artificiale [46], [47].

- Faza 4 (2015-sot) DL (Deep Learning) dhe GAI (Generative AI)

Rrjetet neurale të thella dhe modelet Transformer mundësuan avancime të mëdha në NLP, vizion kompjuterik dhe sisteme autonome. Modelet LLM (Large Language Modeling) shtuan aftësi të reja në vendimmarrje, duke ngritur njëkohësisht sfida për transparencën dhe etikën [48], [49], [50].

2.1.3. Spektri Autonomisë dhe nivelet e ndërhyrjes njerëzore

Vendimmarrja automatike nuk është një fenomen binar sepse ajo nuk është thjesht "e automatizuar" ose "njerëzore". Studimet e sistemeve njeri-makinë kanë identifikuar që herët ekzistencën e një spektri të autonomisë, i cili shtrihet nga mbështetja e plotë te njeriu deri te funksionimi tërësisht i pavarur i sistemit algoritmik.

Brenda këtij spektri, literatura dallon pesë nivele funksionale, të karakterizuara nga shkalla e delegimit kognitiv ndaj sistemit (Figura 2.1). Në nivelin e parë, sistemi luan rolin e informuesit ai mbledh dhe prezanton të dhëna, ndërsa vendimi mbetet plotësisht në duart e njeriut. Në nivelin e dytë, sistemi sugjeron opsione të renditura, por njeriu ruan autoritetin e zgjedhjes finale. Niveli i tretë shënon pikën e kalimit ku sistemi fillon të veprojë vetë, duke i lënë njeriut vetëm mundësinë e ndërprerjes brenda një dritareje të kufizuar kohore. Në nivelin e katërt, sistemi ekzekuton vendimin dhe njofton njeriun pas faktit, duke e reduktuar rolin e tij në mbikëqyrje pasive. Së fundi, niveli i pestë përfaqëson autonominë e plotë sistemi vendos dhe vepron pa asnjë njoftim apo ndërhyrje njerëzore. Nivelet e pëmbledhura paraqiten te figura

Niveli	Emërtimi	Roli i algoritmit	Ndërhyrja njerëzore	Shembull
1	Informues (Sugjerues pasiv)	Mbledh dhe prezanton të dhëna; nuk propozon asnjë veprim	E plotë — njeriu merr çdo vendim	Dashboard, raport analitik
2	Sugjerues (Rekomandues)	Propozon opsione të renditura; njeriu zgjedh njërin prej tyre	E lartë — aprovim manual i çdo opsioni	Netflix, GPS, filtrat e postës
3	Ekzekutues i kushtëzuar	Fillon të veprojë vetë; njeriu mund ta ndalë brenda dritares kohore	E mesme — veto e mundshme, jo e detyrueshme	Email me anulim 10 sek, filtrim spam
4	Ekzekutues i pavarur	Ekzekuton vendimin; njofton njeriun vetëm pas faktit	E ulët — mbikëqyrje pasive, pa ndërhyrje direkte	Bilokimi i kartës bankare, alarm anti-mashtrimi
5	Plotësisht autonom	Vendos dhe vepron pa asnjë njoftim apo ndërhyrje	Zero — sistemi operon tërësisht i pavarur	HFT financiar, agjentë LLM autonom

Kontroll njerëzor Autonomi e plotë

Figura 2.1 Nivelet e automatizimit të vendimmarrjes dhe roli i algoritmit në ndërhyrje

2.1.4. Kuadri Rregullator si Pikë Referimi Konceptuale

Kuadri rregullator përbën një element thelbësor në ndërtimin konceptual të sistemeve të vendimmarrjes automatike, duke shërbyer si një strukturë udhëzuese që përcakton kufijtë,

përgjegjësitë dhe standardet etike të përdorimit të inteligjencës artificiale. Në kontekstin e zhvillimit të sistemeve të bazuara në Inteligjencë Artificiale, siç janë DT (Digital Twin) për menaxhimin e cilësisë së ajrit të brendshëm apo modelet e analizës së sentimentit, kuadri rregullator siguron që proceset automatike të jenë transparente, të përgjegjshme dhe të drejtëpërdrejta për përdoruesin.

Një nga referencat kryesore në këtë drejtim është Rregullorja e Përgjithshme për Mbrojtjen e të Dhënave (GDPR) e Bashkimit European, e cila përfshin dispozita specifike për vendimmarrjen automatike dhe profilizimin (Neni 22). Kjo rregullore thekson të drejtën e individëve për të mos iu nënshtruar vendimeve që bazohen vetëm në përpunim automatik, si dhe nevojën për ndërhyrje njerëzore dhe shpjegueshmëri [51]. Kjo lidhet drejtpërdrejt me konceptin e XAI (Explainable AI), i cili synon të bëjë të kuptueshme dhe të interpretuara vendimet e modeleve komplekse.

Përtej GDPR, iniciativa të tjera si AI Act i Bashkimit European propozojnë një qasje të bazuar në risk, duke klasifikuar sistemet e Inteligjencës Artificiale sipas ndikimit të tyre dhe duke vendosur kërkesa më të rrepta për sistemet me risk të lartë. Në këtë kuadër, aplikimet që lidhen me shëndetin, mjedisin apo vendimmarrjen automatike për përdoruesit konsiderohen kritike dhe kërkojnë auditueshmëri, transparencë dhe kontroll njerëzor [52].

Në aspektin konceptual, kuadri rregullator nuk duhet parë vetëm si një mekanizëm kufizues, por si një instrument që orienton projektimin e sistemeve inteligjente drejt besueshmërisë dhe pranueshmërisë sociale. Në rastin e integritetit të Digital Twin me modele parashikuese dhe module vendimmarrjeje, respektimi i këtyre standardeve siguron që vendimet automatike (p.sh., aktivizimi i një dehumidifier-i bazuar në parashikim) të jenë të justifikuara, të interpretuara dhe të kontrollueshme nga përdoruesi.

2.1.5. Vendimmarrja automatike si sistem socio-teknik

Vendimmarrja automatike, ADM, përfaqëson një evolucion të mëtejshëm të automatizimit të funksioneve njohëse, ku vendimet ose parashikimet që i zëvendësojnë ato delegohen te sistemet algoritmike dhe përdoren për veprime me pasoja reale [53], [54]. Ndryshe nga sistemet tradicionale të mbështetjes së vendimit, DSS (Decision Support System), ku teknologjia ka rol këshillues, ADM karakterizohet nga transferimi i drejtpërdrejtë i funksioneve vendimmarrëse drejt inteligjencës artificiale [55]

Në këtë kuadër, ADM duhet kuptuar si një sistem socio-teknik, ku ndërthuren algoritmet, të dhënat dhe konteksti institucional, duke kërkuar jo vetëm optimizim të performancës, por edhe transparencë, mbikëqyrje dhe llogaridhënie. Me zhvillimin e ML (Machine Learning) dhe data-driven AI, këto sisteme janë bërë më adaptive dhe të shkallëzueshme, duke gjetur aplikim në fusha të ndryshme si financa, shëndetësia, administrata publike dhe sistemet inteligjente të mjedisit

2.2. Modelet teorike të vendimarrjes

Modelet teorike të vendimarrjes përfaqësojnë një bazë konceptuale për të kuptuar mënyrën se si individët dhe sistemet marrin vendime në kushte të ndryshme pasigurie dhe kompleksiteti. Këto modele ndahen në qasje normative, deskriptive dhe preskriptive, duke ofruar perspektiva të ndryshme mbi racionalitetin, sjelljen njerëzore dhe optimizimin e

vendimeve. Në këtë nënkapitull trajtohen teoritë kryesore të vendimmarrjes, duke përfshirë racionalitetin e kufizuar, teorinë e perspektivës dhe teorinë e utilitetit të pritur, si dhe roli i tyre në zhvillimin e sistemeve moderne të inteligjencës artificiale dhe vendimmarrjes automatike.

2.2.1. Paradigmat normative, deskriptive dhe preskriptive

Modelet normative fokusohen në mënyrën se si duhet të merren vendimet në kushte ideale racionaliteti, duke synuar optimizimin e rezultateve. Modelet deskriptive analizojnë mënyrën reale se si individët marrin vendime, duke reflektuar kufizimet njohëse dhe kontekstuale. Ndërsa modelet preskriptive synojnë të kombinojnë të dyja qasjet, duke ofruar metoda praktike për përmirësimin e vendimmarrjes, shpesh të integruara me kuadrin e zhvillimit të modeleve ML (Machine Learning) dhe proceset e NDP (New Product Development).

2.2.2. Teoria e Racionaliteti të Kufizuar

Herbert Simon, një studiues amerikan dhe si një nga figurat kryesore në studimin e vendimmarrjes, argumenton se individët nuk janë plotësisht racionalë, por operojnë nën kufizime informacioni, kohe dhe kapaciteti njohës [56]. Koncepti i Racionalitetit të Kufizuar shpjegon pse vendimmarrja reale shpesh synon “zgjidhje të mjaftueshme” (satisficing) dhe jo optimale, duke krijuar bazën teorike për automatizimin e vendimmarrjes, ADM (Automated Decision Making) dhe përdorimin e sistemeve inteligjente që kompensojnë këto kufizime.

2.2.3. Teoria e Perspektivës, heuristikat dhe bias-et

Kahneman dhe Tversky prezantojnë teorinë e perspektivës, duke treguar se vendimet ndikohen nga heuristikat (rregulla të shpejta mendimi) dhe bias-et (paragjykime sistematike) njohëse, pra njerëzit nuk janë plotësisht racional [57]. Këto fenomene nuk kufizohen vetëm te njerëzit, por studimet e fundit tregojnë se edhe modelet e mëdha gjuhësore, LLM (Large Language Modeling) mund të shfaqin bias-e të ngjashme, duke e bërë të rëndësishme integrimin e mekanizmave të kontrollit dhe shpjegueshmërisë në sistemet IA.

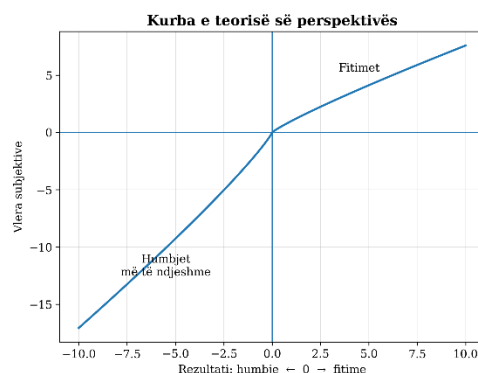


Figura 2.2 Kurba e teorisë së perspektivës me vlerën subjektive të fitimeve dhe humbjeve

Figura 2.2 paraqet kurben e teorisë së perspektivës [57] e cila tregon se si individët perceptojnë fitimet dhe humbjet. Boshti horizontal paraqet rezultatin (humbje ose fitim), ndërsa boshti vertikal paraqet vlerën subjektive. Për fitimet, kurba rritet me ritëm më të ngadaltë, që tregon se individët bëhen më pak të ndjeshëm ndaj fitimeve të mëdha dhe për humbjet, kurba është më e theksuar duke theksuar se ato perceptohen më fort se fitimet. Rezultati është se individët janë më të ndjeshëm ndaj humbjeve sesa ndaj fitimeve, duke reflektuar aversionin ndaj humbjes sipas teorisë së perspektivës.

2.2.4. Teoria e Utilitetit të Pritur

Teoria e utilitetit të pritur, SEU (Subjective Expected Utility), e prezantuar nga Von Neumann dhe Morgenstern (1944), e shpjegon vendimarrjen si një proces në të cilën individët zgjedhin alternativën që sjell përfitimin më të lartë të pritshëm. Kjo teori konsiderohet një model normativ themelor, pasi përshkruan mënyrën sesi, në aspektin teorik, duhet të merren vendimet racionale (pra tregon si “DUHET” të merren vendimet, Figura 2.3 [58]. Megjithatë, kjo qasje është sfiduar nga Simon dhe Kahneman, të cilët tregojnë se në praktikë individët devijojnë nga racionaliteti perfekt, duke theksuar nevojën për modele më realiste dhe adaptive në IA.

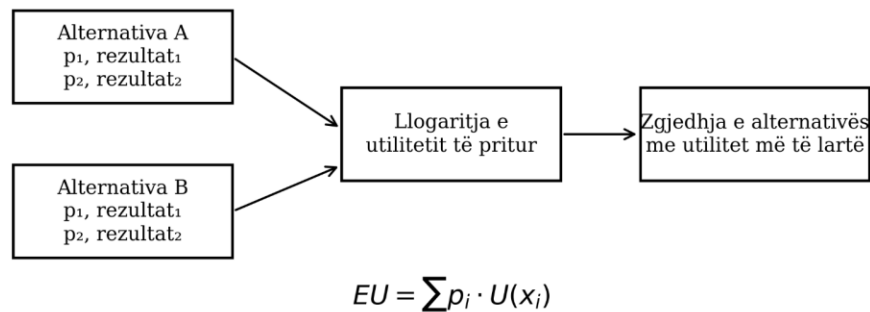


Figura 2.3 Vendimarrja sipas SEU me alternativa për maksimizimin e Utilitetit të Pritur

Figura 2.3 paraqet formulën e Teorisë së Utilitetit, ku :

- p_i është probabiliteti që ndodh rezultati x_i
- x_i është rezultati (p.sh. fitim, humbje)
- $U(x_i)$ është utiliteti (sa vlerë ka për ty ai rezultat)
- EU është utiliteti total i pritur

Pra shumëzohet çdo rezultat me probabilitetin e tij dhe më pas mblidhen të gjitha. Megjithatë, kjo formulë supozon racionalitet të plotë dhe njohje të saktë të probabiliteteve, gjë që shpesh nuk përmbushet në vendimarrjen reale.

2.2.5. Sintetizimi modeleve teorike dhe roli tyre në IA

Modelet teorike të vendimarrjes, duke përfshirë qasjen normative SEU (Subject Expected Utility), qasjet deskriptive [56], [57] dhe qasjet preskriptive, ofrojnë një kuadër të plotë për kuptimin dhe përmirësimin e proceseve vendimmarrëse. Ndërsa modelet normative fokusohen në optimizimin racional, modelet deskriptive reflektojnë kufizimet reale dhe sjelljen njerëzore, duke theksuar devijimet nga racionaliteti perfekt. Në

kontekstin e inteligjencës artificiale, këto modele integrohen për të ndërtuar sisteme të vendimmarrjes automatik, ku algoritmet bazohen në optimizim, por gjithashtu marrin në konsideratë pasigurinë, bias-et dhe nevojën për shpjegueshmëri.

Në arkitekturën e DT (Digital Twin), sensorët e cilësisë së ajrit ofrojnë të dhëna të cilat, megjithëse të dobishme, nuk janë gjithmonë të plota apo perfekte, duke reflektuar konceptin e racionalitetit të kufizuar të Simon. Këto të dhëna përpunohen nga modeli i Machine Learning, i cili gjeneron parashikime dhe probabilitete, duke u lidhur me logjikën e teorisë së utilitetit të pritur, SEU. Më pas, moduli i vendimmarrjes përdor këto rezultate për të zgjedhur veprimin optimal, si për shembull ndezjen ose fikjen e dehumidifier-it, duke operacionalizuar një sistem të vendimmarrjes automatike. Në këtë proces, integrimi i mekanizmave të shpjegueshmërisë është thelbësor për të analizuar dhe kontrolluar bias-et dhe besueshmërinë e vendimeve, në përputhje me qasjet e Kahneman dhe Tversky mbi heuristikat dhe devijimet nga racionaliteti perfekt.

2.3. Nga teoria në algoritëm: teknikat e IA në vendimarrje automatike

Në seksionin 2.2 janë analizuar modele teorike të rëndësishme të vendimmarrjes, veçanërisht teoria e racionalitetit të kufizuar dhe teoria e perspektivës, si dhe teoria e utilitetit të pritur, ofrojnë kuadrin normativ dhe deskriptiv të vendimmarrjes njerëzore. Teknikat e inteligjencës artificiale vijnë si përgjigje praktike ndaj kufizimeve të këtij kuadri: ato kompensojnë kapacitetin njohës të kufizuar, trajtojnë pasigurinë e të dhënave dhe afrojnë vendimmarrjen reale drejt asaj normative, duke e bërë këtë në mënyrë sistematike, të shkallëzueshme dhe, gjithnjë e më shumë, të shpjegueshme. Megjithatë, në literaturën shkencore evidentohet se e njëjta teknikë e inteligjencës artificiale nuk luan gjithmonë të njëjtin rol në procesin e vendimmarrjes. Sipas [53], efekti i një teknike varet nga "pushteti" që i jepet asaj brenda ciklit të vendimit: e njëjta teknikë mund të jetë thjesht informuese në nivelin 1, rekomanduese në nivelin 2 ose plotësisht autonome në nivelin 3. Ky dallim është thelbësor për klasifikimin dhe vlerësimin e teknikave të IA brenda sistemeve ADM (Automated Decision Making). Bazuar në literaturën shkencore dhe tre rastet studimore të këtij disertacioni, chatbot hibrid me XAI (eXplainable AI) në hoteleri, analiza e sentimentit nga e folura në gjuhën shqipe dhe Digital Twin për menaxhimin e cilësisë së ajrit, teknikat e inteligjencës artificiale për vendimarrje automatike organizohen në pesë kategori teorike: Modele Probabilistike, ML (Machine Learning) i orientuar drejt vendimeve, DL (Deep Learning), modele gjuhësore të mëdha dhe Inteligjenca Artificiale e Shpjegueshme si shtresë horizontale e detyrueshme mbi të gjitha kategoritë.

2.3.1. Modelet probabilistike dhe vendimarrja nën pasiguri

Modelet probabilistike përfaqësojnë ndër qasjet teorikisht më të qëndrueshme për vendimarrje automatike, veçanërisht në mjedise të karakterizuara nga pasiguri dhe informacion i pjesshëm. Ky është kushti i zakonshëm i sistemeve reale psh sensorët IoT prodhojnë vlera të cunguara dhe përjashtime, apo rasti vlerësimet të audios tek e cila mund

të ketë ndikime emocionale ose zhurma të tjera. Në të gjitha këto kontekste, modelet probabilistike nuk prodhojnë thjesht një klasifikim binar, por modelojnë shprehimisht pasigurinë përmes shpërndarjeve probabilistike [59].

Një nga përfaqësimet më të rëndësishme të kësaj qasjeje janë Rrjetet Bayesiane (Bayesian Networks, të cilat përbëhen nga grafe të orientuar pa cikle, ku çdo njeje përfaqëson një variabël të rastësishëm, ndërsa harkët shprehin varësi kondicionuese midis tyre. Vlera kryesore e këtyre modeleve në kontekstin e vendimmarrjes automatike qëndron në aftësinë e inferencës probabilistike: duke njohur gjendjen e disa variablave, sistemi mund të vlerësojë probabilitetin e variablave të tjerë të panjohur. Ky mekanizëm është i lidhur ngushtë me teorinë e utilitetit të pritur (SEU) [60], ku vendimi optimal merret duke maksimizuar utilitetin e pritur, por në këtë rast probabilitetet nuk janë të paracaktuara nga eksperti, por mësohen nga të dhënat. Në kontekstin e Digital Twin (Kapitulli VI), kjo qasje mund të përdoret për të vlerësuar probabilitetin e tejkalimit të pragut të lagështirës dhe për të mbështetur vendimin për aktivizimin e thithësit të lagështirës.

Një tjetër klasë e rëndësishme janë HMM (Hidden Markov Model), të cilat janë modele probabilistike sekuenciale (në të cilat konsiderohet edhe koha) ku sistemi përshkruhet nga gjendje të fshehura, të cilat nuk mund të vëzhgohen drejtpërdrejt, por inferohen përmes të dhënave të observueshme. Këto modele janë veçanërisht relevante për analizën e të folurës, ku gjendja emocionale e folësit konsiderohet si variabël e fshehur, ndërsa sinjali akustik përfaqëson vëzhgimin. Edhe pse në Kapitullin V janë përdorur metoda të mësimin të thellë, HMM-të përbëjnë një bazë të rëndësishme teorike për zhvillimin e sistemeve të njohjes automatike të të folurës dhe analizës së sentimentit. Nga perspektiva e vendimmarrjes automatike, modelet probabilistike janë veçanërisht të përshtatshme për nivelet 2–3 të automatizimit sipas klasifikimit të [53], pasi ato jo vetëm që japin një parashikim, por edhe një masë të pasigurisë së tij. Kjo lejon sistemin të vendosë nëse duhet të veprojë në mënyrë autonome apo të kërkojë ndërhyrje njerëzore. Një mekanizëm i tillë, që mund të përshkruhet si “vendimmarrje mbi vetë vendimin”, përbën një karakteristikë thelbësore të sistemeve të besueshme të vendimmarrjes automatike.

2.3.2. ML (Machine Learning) i orientuar drejt vendimeve

Në literaturën bashkëkohore të inteligjencës artificiale, një dallim i rëndësishëm konceptual bëhet ndërmjet modeleve të machine learning të orientuara drejt analizës dhe atyre të orientuara drejt vendimmarrjes. Modelet analitike synojnë kryesisht nxjerrjen e njohurive nga të dhënat përmes parashikimit, klasifikimit apo zbulimit të strukturave të fshehura [46], [61], ndërsa modelet e orientuara drejt vendimit integrojnë këto rezultate në procese që gjenerojnë veprime konkrete, shpesh në mënyrë automatike [54], [55]. Ky tranzicion nga “parashikim” në “vendim” përfaqëson një hap thelbësor në zhvillimin e sistemeve të vendimmarrjes automatike (ADM), ku algoritmet nuk kufizohen vetëm në vlerësimin e gjendjeve të mundshme, por kontribuojnë drejtpërdrejt në zgjedhjen e alternativave dhe ekzekutimin e veprimeve. Në këtë kuptim, teknikat e machine learning shërbejnë si bazë për ndërtimin e funksioneve vendimmarrëse, duke e zhvendosur fokusin nga analiza e të dhënave drejt operacionalizimit të vendimeve në sisteme inteligjente. Machine Learning (ML) klasik përbën një nga shtyllat kryesore të sistemeve të vendimmarrjes automatike (ADM) në implementimet praktike. Ndryshe nga modelet probabilistike, të cilat modelojnë në mënyrë eksplicite pasigurinë, modelet ML mësojnë rregullat vendimmarrëse

drejtpërdrejt nga të dhënat e etiketuara. Si rezultat, ato ndërtojnë funksione vendimi të optimizuara ndaj objektivave të matshëm, si saktësia apo gabimi minimal [46]. Për të kuptuar më mirë rolin e teknikave të ML në sistemet e vendimmarrjes automatike, ato klasifikohen sipas funksionalitetit të tyre në analizë, parashikim dhe vendim, siç paraqitet në Tabelën 2.1

Tabela 2.1 Klasifikimi i teknikave të ML sipas funksionalitetit të tyre

Teknika/ modeli	Analizë	Parashikim	Vendim	Përshkrim (çfarë bëjnë)
LR -Linear Regression	✓	✓	✗	Model linear për parashikimin e vlerave numerike
Arima- Autoregressive integrated moving average	✗	✓	✗	Model statistikor për parashikim të serive kohore
K-means	✓	✗	✗	Grupim i të dhënave në cluster pa etiketa
Pca - Principal Component Analysis	✓	✗	✗	Redukton dimensionet e të dhënave duke ruajtur informacionin kryesor
Clustering	✓	✗	✗	Zbulon struktura dhe ngjashmëri në të dhëna
DT-Decision Trees	✓	✓	✓	Model me rregulla “if-then” që mund të japë vendime direkte
Random Forest	✓	✓	✓	Kombinon shumë pemë për vendime më të qëndrueshme
XGBoost	✓	✓	✓	Model boosting që optimizon gabimet dhe rrit saktësinë
LightGBM	✓	✓	✓	Version efikas i boosting për vendimmarrje të shpejtë
SVM- Support Vector Machine	✓	✓	✓	Model klasifikimi që ndan të dhënat në kategori
Naive Bayes	✓	✓	✓	Model probabilistik për klasifikim dhe vendim
LSTM / GRU	✓	✓	✓	Modele sekuenciale për të dhëna kohore dhe vendim dinamik
Reinforcement Learning	✗	✗	✓	Model që mëson veprime optimale bazuar në shpërblim
MDP -Markov Decision Process	✗	✗	✓	Model matematikor për vendimmarrje në kohë
Rule-based system	✗	✗	✓	Vendimmarrje bazuar në rregulla fikse “if-then”

Pemët e vendimmarrjes (Decision Trees) përfaqësojnë një nga format më intuitive të vendimmarrjes automatike, duke strukturuar procesin si një sekuencë pyetjesh binare mbi karakteristikat hyrëse. Çdo rrugë nga rrënja deri në një nyje fundore (gjethe) përfaqëson një rregull të qartë vendimmarrës. Avantazhi kryesor i këtyre modeleve qëndron në interpretueshmërinë e tyre të natyrshme, çka i bën veçanërisht të përshtatshme për sisteme ku transparencia dhe auditueshmëria janë kritike, si në kontekstin e GDPR (General Data Protection Regulation) dhe AI Act [62]. Megjithatë, ato janë të prirura ndaj overfitting kur rriten pa kufizime të kontrolluara.

Për të adresuar këtë kufizim, Random Forest përdor një qasje ansambli, duke kombinuar shumë pemë vendimesh të trajnuara mbi nëngrupe të ndryshme të të dhënave dhe karakteristikave [63]. Vendimi final merret përmes votimit të shumicës (për klasifikim) ose mesatares (për regresion), duke rritur ndjeshëm robustësinë dhe performancën. Në

Kapitullin IV, Random Forest është përdorur si komponent vendimmarrës në një chatbot hibrid, duke arritur një saktësi bazë prej 0.82 në klasifikimin e ashpërsisë së rasteve të klientëve (severity_high). Megjithatë, ky përmirësim në performancë vjen me koston e reduktimit të interpretueshmërisë, pasi vendimi final është rezultat i shumë modeleve individuale dhe kërkon metoda shtesë si SHAP për shpjegim.

Gradient Boosting, dhe veçanërisht implementimet moderne si XGBoost (Extreme Gradient Boosting) dhe LightGBM (Light Gradient Boosting Machine), ndërtojnë modele ansambli në mënyrë sekuenciale, ku çdo model i ri synon të korrigjojë gabimet e modelit paraardhës [64], [65]. Kjo qasje iterative mund të interpretohet teorikisht në përputhje me konceptin e racionalitetit të kufizuar të Simon, ku vendimi optimal nuk arrihet menjëherë, por përmirësohet gradualisht përmes rafinimeve të njëpasnjëshme. Në Kapitullin VI, XGBoost ka demonstruar performancë shumë të lartë ($F1 = 0.990$, $Recall = 0.990$) në vendimmarrjen për aktivizimin e dehumidifier-it, duke e konfirmuar si një teknikë shumë efektive për seri kohore me sjellje jolineare.

Support Vector Machines (SVM) përfaqësojnë një tjetër qasje të rëndësishme, e cila synon të gjejë hiperplanin optimal që ndan klasat duke maksimizuar distancën (marginin) midis tyre. Këto modele janë veçanërisht të qëndrueshme në kushte ku numri i karakteristikave është i madh dhe dataset-i është relativisht i vogël [66]. Edhe pse nuk janë përdorur drejtpërdrejt në rastet studimore të këtij disertacioni, SVM-të mbeten një referencë teorike e rëndësishme, sidomos në problemet e klasifikimit të sinjaleve komplekse, si analiza e emocioneve nga të folura.

2.3.3. DL (Deep Learning) dhe vendimmarrja e bazuar në të dhëna të pastruara

Në kuadër të zhvillimit të sistemeve moderne të vendimmarrjes automatike, teknikat e Deep Learning kanë marrë një rol gjithnjë e më të rëndësishëm falë kapacitetit për të identifikuar modele të ndërlikuara dhe varësi jo-lineare në të dhëna. Në dallim nga qasjet klasike të të mësuarit makinerik, të cilat zakonisht mbështeten në përcaktimin manual të veçorive (feature engineering), modelet e DL mundësojnë nxjerrjen automatike të karakteristikave përmes arkitekturave me shumë shtresa. Kjo i bën ato më të përshtatshme për përpunimin e të dhënave të pastruara, si audio, tekst dhe imazhe, por gjithashtu të aplikueshme edhe në analiza të serive kohore dhe të dhënave të strukturuar. Megjithatë, përdorimi i tyre shoqërohet me sfida të rëndësishme, si kompleksiteti i lartë, kërkesat për volum të madh të dhënash dhe mungesa e interpretueshmërisë, çka e bën të domosdoshëm integrimin e teknikave të Explainable AI në kontekstin e vendimmarrjes. DL (Deep Learning) zgjeron aftësinë e ML klasike drejt të dhënave të pastruara si zëri, teksti dhe sinjalet kohore, të cilat nuk mund të shtjellohen me karakteristika të dizajnuara manualisht. Rrjetet neurale të thella mësojnë hierarki përfaqësimesh: shtresat e hershme kapin karakteristika të nivelit të ulët (kufij, tone), ndërsa shtresat e thella kapin koncepte semantike abstrakte [49]. Kjo aftësi "end-to-end learning" duke filluar nga sinjali bruto deri te vendimi, e bën DL-në thelbësore për ADM në domaine ku inxhinieria manuale e karakteristikave është e paplotë ose e pamundur. Tabela 2.2 paraqet teknikat DL në ndihmesë të realizimit të vendimarrjes.

Tabela 2.2 Teknikat DL (Deep Learning) për të dhëna dhe vendimmarrje

Teknika	Të dhëna të strukturuar	Të dhëna të pastrukturuar	Përdoret për vendimmarrje	Përshkrimi
MLP (FeedForward NN)	✓	△ (e kufizuar)	✓	Rrjet nervor bazik që modelon marrëdhënie jo-lineare në të dhëna tabelare; përdoret për klasifikim dhe regresion
CNN (Convolutional Neural Networks)	✗	✓✓	✓ (indirekt)	Nxjerr automatikisht karakteristika nga imazhe, audio (spectrogram), dhe tekst; shumë efektiv për perceptim
RNN (Recurrent Neural Networks)	✓	✓	✓	Modelon varësi sekuenciale në të dhëna kohore ose tekst
LSTM	✓✓	✓	✓✓	Variant i avancuar i rnn që kap varësi afatgjata në time series dhe sekuenca
GRU	✓✓	✓	✓✓	Version më i thjeshtë dhe më efikas i LSTM për të dhëna sekuenciale
Transformers (BERT, GPT, etj.)	✗	✓✓✓	✓✓	Arkitekturë moderne me mekanizëm attention; shumë e fuqishme për tekst, audio dhe multimodal
Autoencoders	✓	✓	✗	Përdoren për reduktim dimension, denoising dhe nxjerrje karakteristikash
Variational Autoencoders (VAE)	✓	✓	✗	Gjenerojnë të dhëna të reja duke mësuar shpërndarjen probabilistike
GANs (Generative Adversarial Networks)	✗	✓✓	✗	Gjenerojnë të dhëna realiste (imazhe, audio); përdoren për simulim dhe augmentim
Deep Belief Networks (DBN)	✓	✓	✗	Modele probabilistike të thella për nxjerrje karakteristikash (më pak të përdorura sot)
Deep reinforcement learning (DRL)	✓✓	✓✓	✓✓✓	Model që mëson politika vendimmarrjeje nga ndërveprimi me mjedisin
Graph Neural Networks (GNN)	✓ (strukturë grafi)	△	✓	Përdoren për të dhëna në formë grafi (rrjete sociale, sensorë të lidhur)
Temporal Convolutional Networks (TCN)	✓✓	✓	✓	Alternativë ndaj rnn për time series, me performancë të mirë dhe stabilitet
Attention-based RNN/LSTM	✓✓	✓✓	✓✓	Përmirëson interpretueshmërinë duke fokusuar pjesët më të rëndësishme të input-it
Modele Hibride (CNN+LSTM, etj.)	✓✓	✓✓✓	✓✓	Kombinojnë modele për të kapur karakteristika hapësinore dhe kohore njëkohësisht

Rrjeti Nervor Konvolucional, CNN (Convolutional Neural Network) kapin modelet lokale dhe invariancën hapësinore nëpërmjet filtrave të konvolucionit. CNN-të, të krijuara fillimisht për përpunimin e imazheve, janë përshtatur me sukses edhe për analizën audio përmes trajtimit të spektrogrameve Mel si forma vizuale të sinjaleve akustike. Je

mundëson nxjerrjen automatike të karakteristikave akustike relevante për njohjen e emocioneve pa ndërhyrje manuale të inxhinierit.

LSTM (Long Short-Term Memory) dhe GRU (Gated Recurrent Unit) janë arkitektura rekurrente të dizajnuara për të kapur varësi afatgjata në sekuenca kohore. Porta e harrimit (forget gate) e LSTM-it mundëson selektimin selektiv të informacionit që mbahet ose hidhet, mekanizëm analogjik me aftësinë njerëzore të fokusimit të vëmendjes [67]. Në Kapitullin VI, LSTM dhe GRU janë krahasuar me modelet e bazuara në pemë për parashikimin e lagështirës: megjithëse arkitekturalisht të avancuara, nuk i tejkaluan XGBoost-in dhe Regresionin Linear, gjetje e rëndësishme teorike që tregon se kompleksiteti arkitekturor nuk garanton domosdoshmërisht performancë superiore, veçanërisht kur dataset-i ka strukturë relativisht lineare dhe volum të moderuar.

BiLSTM me shtresë Attention kombinon dy drejtueshmëri sekuenciale me mekanizëm vëmendjeje selektive: modeli lexon sekuencën nga e majta djathtas dhe anasjelltas, ndërsa Attention-i "fokusohet" te kornizat informativisht më të pasura të sinjalit. Matematikisht, Attention llogarit një shpërndarje peshash mbi të gjithë hapat kohore dhe kombinon gjendjet e fshehura në një vektor të ponderuar [68]. Në Kapitullin V, BiLSTM+Attention me karakteristika HuBERT ka arritur performancën maksimale (F1 deri 0.64 për klasën neutrale), duke konfirmuar teorinë se mekanizmat e vëmendjes rrisin interpretueshmërinë e vendimit: mund të identifikohet se cilat pjesë të audios kanë ndikuar më shumë në klasifikimin emocional.

2.3.4. LLM (Large Language Model) dhe vendimmarrja gjeneruese

Vendimmarrja gjeneruese mund të përkufizohet si një qasje në inteligjencën artificiale ku modelet gjeneruese, veçanërisht Large Language Models (LLMs), përdoren për të prodhuar vendime të shoqëruara me arsytetime, rekomandime dhe përmbajtje në gjuhë natyrore, duke mbështetur procesin e vendimmarrjes përtej rezultateve tradicionale numerike ose kategorike. Kjo qasje bazohet në aftësinë e modeleve Transformer për të gjeneruar tekst koherent dhe kontekstual [48], [50] dhe lidhet ngushtë me zhvillimet në Inteligjencën Artificiale Gjenerative dhe DSS (Decision Support Systems) [38], [69].

Modelet e mëdha gjuhësore mund të klasifikohen në disa kategori kryesore bazuar në arkitekturën dhe funksionalitetin e tyre, duke përfshirë modele gjeneruese, encoder-based, encoder-decoder dhe multimodale. Secila kategori ka rol të ndryshëm në përpunimin e të dhënave, duke mbështetur procese si analiza, parashikimi dhe vendimmarrja gjeneruese. Ndërkohë që modelet encoder-based përdoren kryesisht për analiza dhe klasifikim, modelet gjeneruese dhe ato të trajnuara për ndjekjen e udhëzimeve luajnë një rol kyç në mbështetjen e vendimmarrjes, duke ofruar rekomandime dhe arsytetime të strukturuar. Integrimi i këtyre modeleve me burime të jashtme të njohurive, si në rastin e qasjeve RAG (Retrieval Augmented Generation), rrit më tej besueshmërinë dhe saktësinë e vendimeve në sistemet inteligjente. Tabela 2.3 paraqet kategoritë e Modeleve të Mëdha Gjuhësore nëse përfshihen në parashikim, vendimmarrje apo analizë të dhënash.

Tabela 2.3 Kategoritë e LLM (Large Language Model) sipas të dhënave dhe rolit të tyre

Kategoria e LLM	Shembuj	Parashikim	Vendimmarrje	Analizë	Përshkrimi
-----------------	---------	------------	--------------	---------	------------

Generative LLMs	GPT, LLaMA, Claude	✓	✓✓✓	✓✓✓	Gjenerojnë tekst, rekomandime dhe arsyetime; përdoren për vendimmarrje gjeneruese
Encoder-based LLMs	BERT, RoBERTa, DistilBERT	✓✓	✗	✓✓✓	Përdoren për analizë teksti, klasifikim dhe kuptim konteksti
Encoder-Decoder LLMs	T5, BART	✓✓	✓	✓✓✓	Transformojnë tekst në tekst (përkthim, përmbledhje, pyetje-përgjigje)
Multimodal LLMs	GPT-4o, Gemini	✓✓	✓✓✓	✓✓✓	Kombinojnë shumë lloje të dhënash për analizë dhe vendimmarrje komplekse
Instruction-tuned / Chat LLMs	ChatGPT, Claude	✓	✓✓✓	✓✓✓	Optimizuar për dialog, ndjekje udhëzimesh dhe mbështetje vendimmarrjeje
Retrieval-based LLMs (RAG)	GPT + databaza	✓✓	✓✓✓	✓✓✓	Kombinojnë LLM me databaza për vendime të bazuara në njohuri reale
Domain-specific LLMs	BioBERT, FinGPT	✓✓	✓✓✓	✓✓✓	Trajnuar për fusha specifike për vendimmarrje më të saktë
Code LLMs	Codex, CodeLLaMA	✓	✓	✓✓	Përdoren për analizë dhe gjenerim kodi, ndihmojnë në vendimmarrje teknike

Modelet e Mëdha Gjuhësore, LLM (Large Language Models), të bazuara në arkitekturën Transformer [50], përfaqësojnë paradigmen më të re të inteligjencës artificiale për vendimmarrje gjeneruese. Transformer-i zëvendëson rekurrencën e LSTM-it me mekanizëm Self-Attention global: çdo element i sekuencës ndërvepron drejtpërdrejt me çdo element tjetër, duke kapur varësi afatgjata pa kufizimet e rrjeteve rekurrente. Kjo arkitekturë ka mundur shkallëzimin e modeleve deri në qindra miliarda parametra dhe trajnimin mbi korpusë masive tekstuale [48].

Nga perspektiva e vendimmarrjes automatike, LLM-të paraqesin një cilësi të re: aftësinë e gjenerimit të tekstit natyror si përgjigje vendimmarrëse. Ndryshe nga modelet klasike ML që prodhojnë etiketa ose vlera numerike, LLM-të prodhojnë shpjegime gjuhësore, rekomandime të personalizuar dhe dialogje kontekstuale, gjithë që i afron sistemet ADM komunikimit njerëzor. Në Kapitullin IV, GPT-3.5 ka shërbyer si komponenti gjenerues i chatbot-it hibrid: bazuar në vendimin e modulit Random Forest+SHAP, GPT-3.5 ka prodhuar përgjigje profesionale, të personalizuar dhe kontekstualisht koherente ndaj vlerësimeve të klientëve.

Kufizimi kryesor teorik dhe praktik i LLM-ve është saktësisht ajo çfarë i bën të

fuqishme: modelet stokastike probabilitike mbi distribucione masive tekstuale prodhojnë dalje të "kutisë së zezë" (Black Box). Haluçinacionet, prodhimi i informacionit faktikisht të pasaktë me besim të lartë dhe paparashikueshmëria e outputit janë sfida serioze për implementimin e LLM-ve në sisteme ADM me pasoja kritike [70]. Kjo justifikon arkitekturën hibride të Kapitullit IV: LLM-ja është izoluar në rolin e gjenerimit gjuhësor, ndërsa vendimi strategjik i delegohet modelit Random Forest e cila është transparente dhe e auditueshme.

HuBERT (Hidden-Unit BERT) përfaqëson një nënkategori të veçantë të modeleve të vetë-mbikëqyruar (Self-Supervised Learning) të trajnuar mbi audio. Ndryshe nga LLM-të tekstuale, HuBERT nxjerr ngulitje (embeddings) me dimensione 768–1024 drejtpërdrejt nga sinjali audio, duke kapur informacion semantik dhe prozodik pa kërkuar etiketim manual [71]. Në Kapitullin V, HuBERT ka demonstruar superioritet konsistent ndaj MFCC klasike, veçanërisht në kontekstin e gjuhës shqipe si gjuhë me burime të kufizuara, duke konfirmuar se përfaqësimet e mësuara nga të dhëna masive audio transferohen efektivisht edhe në gjuhë dhe dataset me burime të kufizuara.

2.3.5. Inteligjenca Artificiale e Shpjegueshme si shtresë e detyrueshme

Inteligjenca Artificiale e Shpjegueshme, XAI (eXplainable Artificial Intelligence) nuk përbën një teknikë të vetme, por një fushë kërkimore horizontale që përfshin dhe mbështet të gjitha teknikat e inteligjencës artificiale të trajtuara më sipër. Nevoja për XAI lidhet drejtpërdrejt me kufizimet e vendimmarrjes, të njohura edhe në teorinë e Kahneman dhe Tversky [57], sipas së cilës edhe proceset vendimmarrëse, përfshirë ato të modeleve të Inteligjencës Artificiale, mund të ndikohen nga paragjykime sistematike. Në këtë kontekst, modelet si LLM-të mund të trashëgojnë bias-e nga të dhënat e trajnimit, duke e bërë të papranueshëm besimin e pakushtëzuar në rezultatet e tyre, si nga pikëpamja shkencore ashtu edhe etike [72]. Sipas qasjes së DARPA-s (Defense Advanced Research Projects Agency), XAI synon të mundësojë që sistemet e inteligjencës artificiale të shpjegojnë arsyetimin e tyre, të evidentojnë pikat e forta dhe kufizimet, si dhe të ndihmojnë përdoruesin në kuptimin dhe besimin ndaj tyre [73].

Një nga metodat më të rëndësishme të XAI është SHAP (SHapley Additive exPlanations), e cila bazohet në SHAP. SHAP interpreton parashikimet e modeleve të të mësuarit makinerik duke i caktuar secilës karakteristikë një vlerë kontributi të bazuar në koncepte të rrjedhura nga teoria bashkëpunuese e lojërave. E aplikuar në modelet e ML (Machine Learning), SHAP mundëson shpjegimin e çdo parashikimi individual duke treguar se sa ka kontribuar secila karakteristikë në devijimin nga vlera bazë. Vetitë matematikore të saj, si lokaliteti, efikasiteti dhe konsistenca, e bëjnë SHAP një nga metodat më të qëndrueshme teorikisht për shpjegueshmëri post-hoc [74]. Në Kapitullin IV, kjo metodë është përdorur për të gjeneruar shpjegime konkrete, si p.sh.: “Statusi i besnikërisë ka kontribuar 40% në vendimin për përmirësimin e dhomës.”

Një tjetër teknikë e rëndësishme është LIME (Local Interpretable Model-agnostic Explanations), e cila fokusohet në shpjegimin lokal të vendimeve individuale [75]. LIME ndërton modele lineare të thjeshta rreth një instance specifike duke gjeneruar variante të saj (perturbacione) dhe duke analizuar sjelljen e modelit kompleks në atë rajon lokal. Ndryshe nga SHAP, që ofron një pamje më të qëndrueshme dhe globale, LIME është i fokusuar në interpretimin lokal. Kombinimi i këtyre dy metodave, siç është zbatuar në chatbot-in e Kapitullit IV, siguron një nivel të plotë shpjegueshmërie, duke mbuluar si

aspektin global ashtu edhe atë lokal të vendimarrjes.

Përveç metodave post-hoc, ekzistojnë edhe mekanizma të integruar në vetë arkitekturën e modeleve, si Attention, i përdorur në arkitektura si BiLSTM+Attention dhe Transformer. Ky mekanizëm ofron një formë shpjgueshmërie të brendshme (endogjene), pasi peshat e vëmendjes tregojnë se cilat pjesë të input-it kanë ndikuar më shumë në vendim [68]. Në kontekstin e Kapitullit V, kjo qasje është veçanërisht e vlefshme për analizën e të dhënave audio, pasi lejon identifikimin e segmenteve kohore që kontribuojnë më shumë në klasifikimin emocional, duke evidentuar kështu “momentet vendimtare” brenda sinjalit të folurit.

2.3.6. Sintezë: teknikat e IA sipas funksionit në ciklin e Vendimarrjes Automatike

Klasifikimi teorik i teknikave të inteligjencës artificiale bëhet i plotë vetëm kur teknikat shqyrtohen sipas funksionit që luajnë brenda ciklit të vendimarrjes automatike, dhe jo vetëm sipas algoritmit. Duke u bazuar në modelin e [76] procesi i automatizimit përfshin katër funksione kryesore: mbledhja e informacionit, analiza e informacionit, zgjedhja e vendimit dhe zbatimi i veprimit. Funksionalitetet janë përdorur në të tre rastet studimore të këtij disertacioni.

Tabela 2.2 Teknikat e IA sipas funksionit në ciklin vendimarrjes automatike dhe nivelit të shpjgueshmërisë

<i>Teknika</i>	<i>Funksioni në ADM</i>	<i>Roli</i>	<i>Shpjgueshmëria</i>	<i>Rasti studimor</i>
<i>Random Forest</i>	Analiza + Vendim	Klasifikim vendimesh	E mesme (SHAP)	IV – Chatbot
<i>XGBoost</i>	Analiza + Vendim	Parashikim Vendimarrjes	E mesme (SHAP)	VI – DT
<i>LSTM / GRU</i>	Analiza sekuenciale	Parashikim kohore	E ulët	VI – DT
<i>BiLSTM + Attention</i>	Analiza audio	Klasifikim emocional	E mesme (Attention)	V – SER Shqip
<i>HuBERT</i>	Blerja e informacionit	Nxjerrje veçorish	E ulët	V – SER Shqip
<i>GPT-3.5</i>	Zbatimi i veprimit	Gjenerim gjuhësor	E ulët (kuti e zezë)	IV – Chatbot
<i>SHAP / LIME</i>	Shpjegim post-hoc	Transparencë Vendimarrjes	E lartë	IV, VI

Nga tabela 2.2 evidenton tre momente të rëndësishme. Së pari, asnjë teknikë e vetme nuk mbulon të gjithë ciklin realizimit të vendimarrjes automatike: ndërtimi i sistemeve vendimarrëse efektive kërkon kombinim teknikash sipas funksionit. Kjo justifikon arkitekturën hibride të tre rasteve studimore, ku teknika të ndryshme janë integruar në shtresa funksionale. Së dyti, ekziston kompromis i qartë midis fuqisë parashikuese dhe shpjgueshmërisë ku GPT-3.5 dhe LSTM janë më të fuqishme por më pak transparente, ndërsa Random Forest dhe XGBoost balancojnë mirë të dyja dimensionet. Së treti, XAI nuk është opsionale: si kërkesë rregullatore (AI Act, GDPR, Neni 22) dhe si kërkesë shkencore, si bias i modelit, integrimi i teknikave SHAP ose Mekanizmit të Vëmendjes (Attention) është komponent i detyrueshëm i çdo sistemi për një vendimarrje të besueshme.

2.3.7. Qasje të tjera të Inteligjencës Artificiale për vendimmarrje

Përveç teknikave të bazuara në machine learning dhe deep learning, literatura shkencore identifikon edhe një sërë qasjesh të tjera të inteligjencës artificiale që lidhen me vendimmarrjen automatike. Këto përfshijnë modele probabilistike, algoritme optimizimi dhe sisteme të bazuara në rregulla, të cilat kanë kontribuar historikisht në zhvillimin e sistemeve vendimmarrëse.

Modelet probabilistike, si rrjetet Bayesiane, përdoren për modelimin e pasigurisë dhe vendimmarrjen në kushte probabilistike, duke reflektuar qasjen normative të teorisë së utilitetit të pritur. Algoritmet e optimizimit, përfshirë metodat heuristike dhe metaheuristike (p.sh. algoritmet gjenetike), synojnë gjetjen e zgjidhjeve optimale në probleme komplekse me shumë alternativa. Ndërkohë, sistemet e bazuara në rregulla dhe sistemet eksperte përfaqësojnë një nga format më të hershme të automatizimit të vendimmarrjes, ku vendimet merren në bazë të logjikës “nëse-atëherë”.

Në kontekstin modern, këto qasje shpesh kombinohen me modelet e machine learning për të ndërtuar sisteme hibride, të cilat integrojnë njohuritë e ekspertëve me analizën e të dhënave, duke rritur fleksibilitetin dhe besueshmërinë e vendimmarrjes automatike.

Megjithëse këto qasje ofrojnë alternativa të vlefshme për vendimmarrje, ky studim fokusohet në modelet e machine learning dhe deep learning për shkak të përshtatshmërisë së tyre me të dhënat e disponueshme dhe kërkesat e aplikimit në sistemet e analizës së cilësisë së ajrit dhe Digital Twin.

2.4. Algoritmet Optimizues

Algoritmet optimizuese janë metoda matematikore dhe algoritmike që synojnë të identifikojnë zgjidhjen më të mirë të mundshme të një problemi vendimmarrjeje, duke maksimizuar ose minimizuar një funksion objektiv në prani të kufizimeve të caktuara. Ato përdoren për të përzgjedhur alternativën optimale nga një grup opsionesh dhe përbëjnë një komponent thelbësor në procesin e vendimmarrjes, veçanërisht në sistemet e inteligjencës artificiale ku rezultatet parashikuese duhet të shndërrohen në veprime konkrete [77]. Formulimi matematikor i përgjithshëm është:

$$a^* = \arg \max_{a \in \mathcal{A}} \mathbb{E}[U(a, X) \mid \hat{p}(X)] \quad (2.1)$$

a^* : veprimi (alternativa)

U : utiliteti (objektivi)

$\hat{p}(X)$: parashikimi i modelit

Optimizimi zgjedh a^* që jep utilitet maksimal duke u bazuar te parashikimi. Ai kalon disa faza, psh optimizimi gjatë trajnimit modeli bëhet më i saktë por nuk merr vendime. Ndërkohë optimizimi në vendim zgjedh veprimin më të mirë dhe konsiderohet optimizimi real i ADM. Optimizimi në kohë reale ku sistemi mëson dhe përshtatet është niveli më i avancuar.

Tabela 2.4 Klasifikimi i algoritmeve të optimizimit sipas fazës së përdorimit në procesin e vendimmarrjes automatike

<i>Faza e optimizimit</i>	<i>Çfarë optimizohet?</i>	<i>Alternativat / Teknikat</i>	<i>Ku përdoret?</i>
1. Gjatë trajnimit (Model Optimization)	Parametrat e modelit, gabimi (loss)	Gradient Descent, Adam, Hyperparameter tuning, Bayesian Optimization	ML, DL, Audio (Kapitulli V)
2. Pas parashikimit (Decision Optimization)	Veprimi optimal bazuar në output	Threshold optimization, Rule-based, Linear Programming, Multi-objective optimization	Chatbot (IV), Digital Twin (VI)
3. Në kohë reale (Control Optimization)	Vendime dinamike në kohë	Reinforcement Learning, Dynamic Programming, MPC (Model Predictive Control)	Digital Twin (VI)
4. Optimizim strukturor (Model/Feature Optimization)	Zgjedhja e modelit ose inputeve	Feature selection, AutoML, Neural Architecture Search	ML, DL
5. Optimizim multi-objektiv	Balancimi i disa qëllimeve	Pareto Optimization, Weighted sum, Evolutionary algorithms	Digital Twin, sisteme reale

Tabela 2.4 paraqet klasifikimin e algoritmeve të optimizimit në varësi të fazës së përdorimit të tyre në procesin e vendimmarrjes automatike. Siç vihet re, optimizimi nuk kufizohet vetëm në fazën e trajnimit të modeleve, por shtrihet në disa nivele funksionale të sistemit. Në fazën e parë, algoritmet e optimizimit përdoren për përmirësimin e performancës së modelit përmes minimizimit të funksionit të gabimit. Në fazën e dytë, optimizimi shërben për përzgjedhjen e veprimit optimal bazuar në rezultatet e modelit, duke e kthyer parashikimin në vendim konkret. Në fazat më të avancuara, si kontrolli në kohë reale dhe optimizimi multi-objektiv, algoritmet mundësojnë përshtatjen dinamike të vendimeve në mjedise komplekse dhe balancimin e qëllimeve të ndryshme. Kjo ndarje thekson rolin e optimizimit si një komponent kyç në ndërtimin e sistemeve të plota të vendimmarrjes automatike.

KAPITULLI III

Analiza e strukturës së të dhënave

3.1. Qëllimi i analizës së strukturës së të dhënave

Përpara aplikimit të teknikave të inteligjencës artificiale dhe delegimit të vendimmarrjes te sistemet algoritmike, është thelbësore të kuptohet struktura e brendshme e të dhënave mbi të cilat këto sisteme ndërtohen. Analiza e strukturës së të dhënave shërben si hap paraprak për të vlerësuar nëse të dhënat janë të përshtatshme për modelim, parashikim dhe përfundimisht për vendimmarrje automatike. Në kontekstin e këtij studimi, kjo analizë synon të identifikojë karakteristikat kryesore statistikore dhe dinamike të të dhënave, të cilat ndikojnë drejtpërdrejt në zgjedhjen e teknikave të inteligjencës artificiale dhe në besueshmërinë e vendimeve të automatizuara.

3.2. Tipologjia e të dhënave dhe ndikimi në zgjedhjen e teknikave

Në sistemet moderne të bazuara në inteligjencë artificiale, tipologjia e të dhënave përbën një faktor themelor që ndikon në mënyrën se si ndërtohen modelet analitike dhe realizohet vendimmarrja automatike. Literatura shkencore thekson se karakteristikat e të dhënave, si struktura, dimensionimi dhe forma e përfaqësimit, përcaktojnë jo vetëm zgjedhjen e teknikave të inteligjencës artificiale, por edhe kompleksitetin dhe besueshmërinë e vendimeve të prodhuara [78], [79].

Në sistemet e bazuara në inteligjencë artificiale, të dhënat paraqiten në forma dhe struktura të ndryshme, të cilat ndikojnë drejtpërdrejt në mënyrën e përpunimit dhe analizës së tyre. Një pjesë e të dhënave organizohen sipas një strukture të përcaktuar, zakonisht në tabela ose databaza, duke mundësuar përdorim më të drejtpërdrejtë në algoritmet tradicionale të të mësuarit makinerik. Ndërkohë, burime të tjera informacioni, si audio, imazhe apo tekst natyror, nuk ndjekin një organizim të tillë dhe për këtë arsye kërkojnë teknika shtesë të përpunimit për t'u transformuar në forma numerike të përshtatshme për modelet e vendimmarrjes automatike. [80].

Tipologjia e të dhënave ndikon drejtpërdrejt në kompleksitetin e procesit vendimmarrës. Në rastet kur të dhënat janë të strukturuar dhe me dimension të ulët, procesi i vendimmarrjes mund të realizohet me modele më të thjeshta dhe më të interpretuara. Ndërsa në rastet kur të dhënat janë të pastrukturuar ose me dimension të lartë, si në përpunimin e sinjaleve audio apo të dhënave të gjeneruara nga sensorë kompleksë, kërkojnë modele më të avancuara, shpesh të bazuara në DL (Deep learning), të cilat rrisin performancën, por njëkohësisht edhe kompleksitetin dhe sfidat e interpretueshmërisë [81], [82].

Për më tepër, literatura thekson se tipologjia e të dhënave ndikon në mënyrën se si vendimmarrja automatike kalon nga një proces reaktiv në një proces parashikues dhe proaktiv. Në sistemet që përdorin të dhëna të strukturuar kohore, si në rastin e sensorëve të cilësisë së ajrit, vendimet mund të bazohen në modele parashikuese që identifikojnë tendenca dhe ndërhyjnë përpara se të ndodhë një devijim kritik. Në të kundërt, në sistemet që përpunojnë të dhëna të pastrukturuar, si audio, vendimmarrja lidhet më shumë me interpretimin e gjendjeve ose klasifikimin e situatave në kohë reale [83].

Raste reale tregojnë se ndryshimi i tipologjisë së të dhënave sjell ndryshime të ndjeshme

në mënyrën e realizimit të vendimmarrjes automatike. Për shembull, në sistemet e monitorimit mjedisor, ku përdoren të dhëna numerike nga sensorë, vendimmarrja mund të strukturohet në forma të thjeshta si rregulla prag (threshold-based decisions) ose modele parashikuese për kontroll automatik. Ndërsa në sistemet e përpunimit të të folurit, vendimet varen nga nxjerrja e veçorive komplekse dhe nga modele klasifikimi me dimension të lartë, duke e bërë procesin më të ndërlikuar dhe më pak transparent [84], [85].

Në këtë mënyrë, tipologjia e të dhënave nuk përbën vetëm një karakteristike teknike, por një element kyç që ndikon në projektimin e sistemeve të inteligjencës artificiale dhe në mënyrën se si këto sisteme marrin vendime. Kuptimi i saj është i domosdoshëm për ndërtimin e sistemeve të vendimmarrjes automatike që janë të sakta, të besueshme dhe të përshtatura për kontekste reale aplikative.

3.3. Analiza e të dhënave të strukturuar nga sensorët mjedisor

Të dhënat e strukturuar dallohen për informacion të organizuar në mënyrë të qartë, duke bërë të mundur nxjerrjen e njohurive, modeleve dhe vendimeve automatike nga sistemi. Analiza e këtyre të dhënave është një hap cilësor është një hap thelbësor përpara përdorimit të teknikave të inteligjencës artificiale, statistikës dhe deri te vendimmarrja automatike. Disa arsye kryesore janë: cilësia e të dhënave si identifikimi i gabimeve të matjeve, vlerat bosh, vlerat e dyshimta etj, është zbulimi i modeleve dhe marrëdhënieve si trendet, sezonaliteti, korrelacionet apo edhe ndryshimet në kohë. Është përshtatshmëria për modelim si nëse ka mjaftueshëm të dhëna, a është shpërndarja e balancuar, nëse ka linearitet apo jo mes variabla. Modelet AI nuk mund të trajnohen në mënyrë efektive pa normalizim, pa pastrim të dhënash etj. Nëse të dhënat nuk analizohen fillimisht si rezultat modeli mund të japë parashikime të pasakta, mund të mësojë nga gabimet dhe zhurmat dhe në fund të prodhojë vendime jo të besueshme.

3.3.1. Rregullsia kohore dhe cilësia e matjeve

Dataseti bruto është vlerësuar fillimisht mbi frekuencën e matjeve sepse përbën një hap të rëndësishëm në vlerësimin e strukturës kohore. Në sistemet e monitorimit të cilësisë së ajrit të brendshëm, frekuenca e matjeve përcakton rezolucionin kohor dhe ndikon drejtpërdrejt në ndërtimin e modeleve parashikuese. Një frekuencë e rregullt lejon modelimin korrekt të varësive kohore, ndërsa një frekuencë jo uniforme mund të tregojë ndërprerje në regjistrimin e të dhënave ose ndryshime në konfigurimin e sensorit. Dataseti ka 98721 rekorde me frekuenca kohore 5 min si dhe parametrat e matur janë temperatura, lagështira, trysnia e ajrit dhe koha. Periudha e përfshirë është viti 2025.

Duke vlerësuar frekuencën e matjeve është konkluduar që 98664 raste janë matje çdo 5 minutë, pra 99.9 % të rasteve e plotësojnë kriterin fillestar të matjes, ndërkohe vetëm 56 raste (0.1%) janë me intervale të ndryshme që dallojnë si shkëputje e energjisë elektrike ose e lidhjes së internetit të pajisjes matëse. Një informacion më i detajuar përcillet sipas tabelës 3.1 ku sipas muajve janë grupuar datat të cilat ka mungesa matjesh.

Tabela 3.1 Grupimi sipas muajve i datave dhe sasisë së matjeve të munguara

Muaji -Viti	Numri i datave me matje të munguara	Numri i mungesave 5 min
Janar 2025	10	1778
Shkurt 2025	4	9
Mars 2025	8	96
Prill 2025	11	42
Maj 2025	8	1491
Qershor 2025	4	11
Gusht 2025	4	424
Shtator 2025	1	258
Tetor 2025	1	1
Dhjetor 2025	1	13

Në dataset-et reale , sidomos në sensorët IoT (Internet of Things) ekzistojnë gabime në matje, probleme me sensorin, zhurma në të dhëna etj të cilat konsiderohen si anomali por ndonjëherë përfaqësojnë fenomene reale që duhet të merren në konsideratë. Këto anomali ndikojnë në shtrembërimin e mesatares apo të devijimit standard, ndikojnë negativisht në modelet e inteligjencës artificiale, prodhojnë parashikime dhe vendime jo të sakta. Gjithsesi identifikimi i tyre shoqërohet me tre qasje kryesore:

- Identifikimi dhe raportimi
- Zëvendësimi ose kufizimi
- Fshirja nga gabimet apo ndryshimet ekstreme

Anomali të duhet të mos kalojnë një nivel të caktuar, psh $<1-5\%$ ku të konsiderohet si anomali e pranueshme gjithsesi nuk ekziston një numër absolut universal që të vendoset si kufi, sepse kjo varet nga sensori dhe dataset-i.

Tabela 3.2 Vlerësimi i anomalive sipas attributeve të dataset-it

Atributet	Nr rekordeve	Nr anomalive	Përqindja anomalive	Kufiri poshtëm	Kufiri sipërm
Temperatura	98721	0	0	2.5	41.7
Lagështira	98721	0	0	21.5	108.7
Trysnia	98721	2515	2.55	98494.7	101137.1

Tabela 3.2 paraqet një analizë mbi numrin, përqindjen dhe kufirin e poshtëm dhe të sipërm të vlerave që kanë atributet në dataset. Siç vihet re anomalia e vlerësuar është 0.85% si një përqindje mesatare shumë e ulët, prandaj është konsideruar si një anomali normale për periudhën një vjeçare.

Duke qenë se dataset-i përfaqëson seri kohore me frekuencë të lartë matjesh, u analizua varësia kohore ndërmjet vlerave të njëpasnjëshme përmes autokorrelacionit, ACF (Auto Correlation Function). Rezultatet treguan se vlerat aktuale të variablave mjedisore ruajnë lidhje të fortë me matjet paraprake, duke konfirmuar ekzistencën e varësisë temporale në seri kohore. Kjo karakteristikë është veçanërisht e rëndësishme për modelet parashikuese të cilat mbështeten në marrëdhëniet kohore ndërmjet vëzhgimeve të vazhdueshme

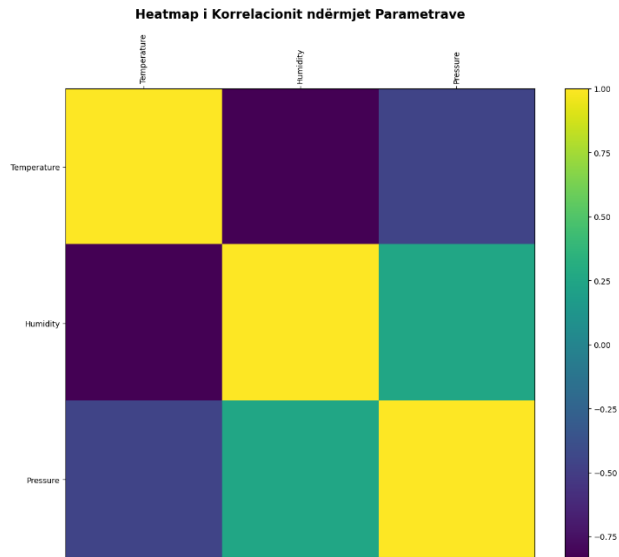


Figura 3.1 Heatmap autokorrelacionit për serinë kohore të lagështirës, temperaturës dhe trysnisë

Figura 3.1 paraqet heatmap-in e korrelacionit ndërmjet parametrave kryesorë mjedisorë të analizuar në dataset, përkatësisht temperaturës, lagështirës relative dhe trysnisë. Rezultatet tregojnë një korrelacion negativ të fortë ndërmjet temperaturës dhe lagështirës relative, duke nënkuptuar se rritja e temperaturës shoqërohet me ulje të lagështirës relative. Kjo sjellje është fizikisht e pritshme dhe përputhet me karakteristikat e ajrit në mjedise të brendshme. Ndërkohë, korrelacioni ndërmjet lagështirës dhe trysnisë rezulton i dobët, duke treguar mungesë të një lidhjeje lineare të fortë ndërmjet këtyre dy parametrave në dataset-in e analizuar. Analiza e korrelacionit ndihmon në kuptimin e marrëdhënieve ndërmjet variablave dhe në vlerësimin e ndikimit të tyre potencial në modelet parashikuese dhe sistemet e vendimmarrjes automatike.

3.3.2. Statistika përshkruese dhe shpërndarja e variablave

Për ta paraqitur edhe më qartë marrëdhënien ndërmjet këtyre parametrave u përdorën grafikët e shpërndarjeve për ekzistencën e linearitetit Figura 3.2.

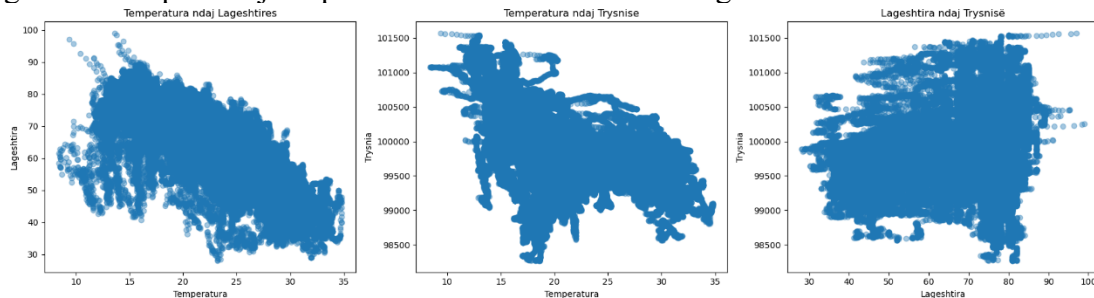


Figura 3.2 Grafiku i shpërndarjes për tre parametrat

Grafikët e shpërndarjes tregojnë marrëdhëniet ndërmjet parametrave kryesorë mjedisorë të dataset-it. Rezultatet evidentojnë një marrëdhënie lineare negative ndërmjet temperaturës dhe lagështirës relative, ku rritja e temperaturës shoqërohet me ulje të lagështirës.

Ndërkohë, marrëdhëniet ndërmjet temperaturës dhe trysnisë, si dhe ndërmjet lagështirës dhe trysnisë, rezultojnë më të dobëta dhe më pak lineare. Këto rezultate tregojnë se dataset-i përmban si varësi lineare ashtu edhe jo-lineare ndërmjet variablave, duke justifikuar përdorimin e modeleve të avancuara të inteligjencës artificiale për analizë dhe parashikim. Figura 3.3 paraqet histogramin e shpërndarjes së vlerave të lagështirës në datasetin e analizuar. Nga figura vërehet se shumica e matjeve janë të përqendruara në intervalin rreth 65%–80%, duke treguar se ambienti ka pasur kryesisht nivele relativisht të larta lagështire gjatë periudhës së monitorimit. Kulmi më i madh i frekuencës shfaqet afërsisht në intervalin 74%–78%, çka tregon se këto vlera janë regjistruar më shpesh krahasuar me intervalet e tjera.

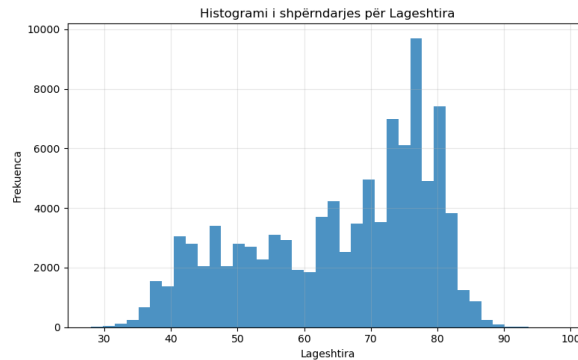


Figura 3.3 Shpërndarja e vlerave të lagështirës sipas frekuencës

Analiza statistikore e variablës së lagështirës (Tabela 3.3) tregoi se në 98,721 matje të realizuar mesatarja është 65.21% dhe devijimi standard 13.5, duke reflektuar variacion të moderuar të kushteve të ambientit. Vlera e shtrembërimit (-0.53) tregon një shpërndarje lehtësisht të anuar majtas, ndërsa kurtoza (-0.93) sugjeron një shpërndarje relativisht të sheshtë me mungesë të anomalive ekstreme. Nga dataseti konfirmohet që 25% e matjeve janë nën 54.2 %, 50 % e matjeve janë në 69 % dhe Në përgjithësi, të dhënat konsiderohen të përshtatshme për analiza statistikore dhe modelim parashikues në kuadër të sistemit Digital Twin dhe vendimmarrjes automatike.

Tabela 3.3 Statistika përshkruese e lagështirës

Mesatarja	65.21
Deviacioni standard	13.5
Minimumi	28
Maximumi	99
25%	54.2
50%	69
75%	76
Shtrembërimi	-0.53
Kurtoza	-0.93

3.3.3. Analiza e pragjeve të lagështirës

Pasi u morr në konsideratë shpërndarja e të dhënave dhe kryesisht e lagështirës ishte e nevojshme vlerësimi i pragut që ndihmon për vendimarrjen si dhe në kuptimin se kur kushtet e ambientit fillojnë të konsiderohen problematike ose të rëndësishme për monitorim, parashikim dhe kontroll automatik.

Matematikisht, pragu më i balancuar mund të identifikohet përmes medianës ose analizës së shpërndarjes mbi/nën pragje të ndryshme. Megjithatë, për shkak se lagështira lidhet me komfortin dhe cilësinë e ajrit të brendshëm, zgjedhja përfundimtare e pragut duhet të kombinojë analizën statistikore me standardet IAQ. Për këtë arsye, në këtë kapitull analizohet shpërndarja e lagështirës sipas pragjeve të ndryshme, ndërsa lidhja me standardet dhe vendimarrjen automatike trajtohet më vonë në Kapitullin VI.

Pragu i llogaritës sipas mesatares është 69% dhe matematikisht konsiderohet prag i balancuar por nga ana e komfortit dhe standardeve IAQ (Indoor Air Quality) ky target nuk është i rregullt. Duke u nisur nga statistikatat dhe përqindjet rezultojnë se 25 % kanë vlera mbi 76% dhe merret në konsideratë për identifikimin e rasteve anormale. Forma më e përshtatshme për përcaktimin e target-it dhe e dobishme për modelet e Inteligjencës Artificiale është testimi i disa pragjeve dhe më pas zgjedhja që krijon raport më të balancuar ndërmjet dy klasave. Duke përfshirë rangun e vlerave nga 45 deri 81 rezultoi që vlera që ruan edhe balancën është përsëri 69. Gjithsesi vlera 69 do jetë perfektja në matematikë por jo vlera e duhur për parashikim. Ndërkohe për t'u përshtatur edhe me nivelet normale të lagështirës në ambiente të brendshme target u vendos një vlerë ndërmjet 60-65 % dhe për pjesën e balancimit të klasës përdoren teknika e peshës së klasave me të cilin modeli i kushton më shumë vëmendje klasës minoritare, teknika SMOTE (Synthetic Minority Oversampling Technique) me anë të të cilës krijohen raste artificiale të klasës së vogël dhe teknika e nën-mostërimit ku tentohet të ulet klasa e madhe. Teknika humbje me fokusim, që përdoret më shumë në teknikat DL (Deep Learning), ku modeli fokusohet më shumë në rastet e vështira etj

3.3.4. Varësia kohore dhe veçoritë me vonesë kohore

Veçoritë me vonesë kohore (lag features) përdoren për të përfshirë informacion nga matjet e mëparshme në modelin parashikues. Këto veçori ndihmojnë modelin të kuptojë varësinë kohore dhe dinamikën e ndryshimit të parametrave të ambientit. Për të përfshirë varësinë kohore të të dhënave, dataset-it iu shtuan veçori me vonesë kohore (lag features), të cilat përfaqësojnë vlerat e parametrave në intervale të mëparshme kohore. Këto veçori i mundësojnë modelit të analizojë trendin dhe zhvillimin kohor të lagështirës, duke përmirësuar aftësinë parashikuese të modeleve të Inteligjencës Artificiale.

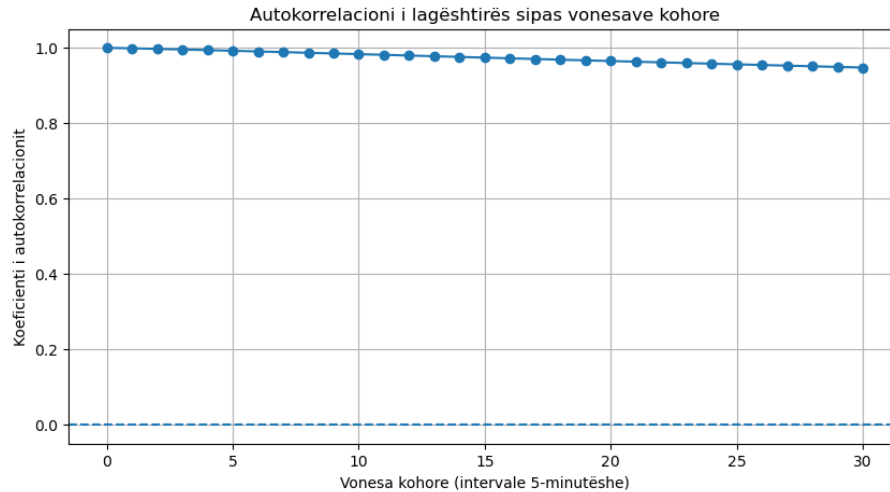


Figura 3.4 Autokorrelacioni i lagështirës relative sipas vonesave kohore në intervale 5-minutëshe.

Figura 3.4 paraqet autokorrelacionin e variablës së lagështirës relative për vonesa kohore të ndryshme (lag features) në intervale 5-minutëshe. Rezultatet tregojnë një autokorrelacion shumë të lartë pozitiv përgjatë të gjitha vonesave të analizuar, ku koeficienti i autokorrelacionit mbetet afër vlerës 1 edhe për lag-et më të largëta. Kjo tregon se lagështira ka varësi të fortë kohore dhe ndryshon gradualisht me kalimin e kohës. Në praktikë, kjo do të thotë se vlerat aktuale të lagështirës ndikohen ndjeshëm nga matjet e mëparshme, duke justifikuar përdorimin e vonesave kohore në ndërtimin e modeleve parashikuese. Gjithashtu, sjellja e qëndrueshme e autokorrelacionit sugjeron se seria kohore ruan vazhdimësi të lartë dhe përmban informacion të vlefshëm historik për parashikimin e kushteve të ardhshme të ambientit të brendshëm.

3.3.5. Stacionariteti dhe sezonaliteti i serisë kohore

Në analizën e serive kohore, stacionariteti përfaqëson aftësinë e serisë për të ruajtur karakteristikat statistikore kryesore, si mesatarja dhe varianca, relativisht të qëndrueshme me kalimin e kohës. Vlerësimi i stacionaritetit është veçanërisht i rëndësishëm për modelet parashikuese tradicionale si ARIMA (AutoRegressive Integrated Moving Average), ndërsa ndikon edhe në performancën e modeleve të inteligjencës artificiale dhe DL (Deep Learning).

Për vlerësimin e stacionaritetit u përdor testi ADF (Augmented Dickey-Fuller). Rezultatet treguan se seria kohore e lagështirës paraqet varësi të fortë kohore dhe ndryshime graduale, ndërkohë që stabiliteti i saj statistikor varet nga periudha kohore e analizuar. Në rastet kur seria nuk rezulton plotësisht stacionare, mund të aplikohen teknika si diferencimi ose transformimet kohore për të përmirësuar përshtatshmërinë për modelim parashikues.

Tabela 3.4 Vlerësimi i stacionaritetit me testin ADF

Nr rreshtave fillestar	Nr rreshtave pas rigrupimit kohor	Vlera llogaritur ADF Statike	P-Value	Vlera Kritike 1%	Vlera Kritike 5%	Vlera Kritike 10%	Rezultati
98721	102844	-9.41257	5.78E-16	-3.430414	-2.861568	-2.566785	-2.566785

Rezultati i tabelës 3.4 paraqet një vlerë ADF prej -9.41 dhe p-value shumë më të vogël se 0.05, duke lejuar refuzimin e hipotezës zero për praninë e rrënjës njësi. Kjo tregon se seria e lagështirës paraqet karakteristika stacionare dhe është e përshtatshme për përdorim në modele parashikuese të serive kohore.

Stacionariteti përbën një kusht themelor për shumë modele klasike të analizës së serive kohore. Për këtë arsye, seritë kohore të përdorura në këtë studim u analizuan për të vlerësuar nëse karakteristikat e tyre statistikore mbeten konstante në kohë. Testet statistikore të stacionaritetit, si Augmented Dickey–Fuller (ADF) dhe KPSS, u përdorën për të identifikuar praninë e trendeve, sezonalitetit dhe ndryshimeve në variancë.

Sezonaliteti në seritë kohore përfaqëson përsëritjen periodike të modeleve ose sjelljeve të caktuara në intervale të rregullta kohore, si p.sh. gjatë ditës, javës, muajit ose stinëve të vitit. Në sistemet e monitorimit të cilësisë së ajrit të brendshëm, sezonaliteti ndihmon në kuptimin e ndryshimeve ciklike të parametrave mjedisorë dhe në identifikimin e sjelljeve normale të ambientit. Analiza e sezonalitetit është veçanërisht e rëndësishme për modelet parashikuese dhe sistemet e vendimmarrjes automatike, sepse mundëson identifikimin e modeleve të përsëritura në kohë dhe përmirëson aftësinë e algoritmeve për të parashikuar vlerat e ardhshme. Për më tepër, njohja e sezonalitetit ndihmon në dallimin ndërmjet ndryshimeve normale periodike dhe anomalive reale në të dhëna.

Figura 3.5 paraqet ndryshimin mesatar të lagështirës relative gjatë orëve të ditës. Rezultatet tregojnë se lagështira është më e lartë gjatë orëve të natës dhe orëve të para të mëngjesit, ndërsa gradualisht ulët gjatë ditës duke arritur vlerat minimale rreth mesditës. Pasdite dhe në mbrëmje vihet re një rritje progresive e lagështirës. Ky model tregon ekzistencën e sezonalitetit ditor në seri kohore dhe sugjeron se parametrat e ambientit ndikohen nga ciklet ditë–natë, temperatura dhe aktivitetet në ambientin e brendshëm. Prania e këtij sezonaliteti është e rëndësishme për modelet parashikuese, pasi informacioni mbi orën e ditës mund të përdoret si veçori hyrëse për përmirësimin e saktësisë së parashikimit.

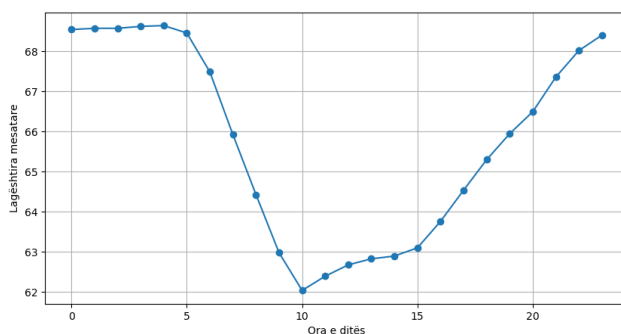


Figura 3.5 Sezonaliteti ditor i lagështirës sipas orës

Figura 3.6 paraqet ndryshimin mesatar mujor të lagështirës relative gjatë vitit. Vërehet një trend i qartë sezonal ku vlerat më të larta të lagështirës regjistrohen gjatë muajve të dimrit, ndërsa vlerat minimale shfaqen gjatë periudhës verore, veçanërisht në muajt qershor–korrik. Kjo sjellje është në përputhje me karakteristikat klimatike dhe ndikimin e temperaturës në lagështirën relative të ambientit të brendshëm. Analiza e sezonalitetit mujor konfirmon ekzistencën e cikleve afatgjata në seri kohore dhe ndihmon modelet parashikuese të kuptojnë ndryshimet sezonale të kushteve mjedisore.

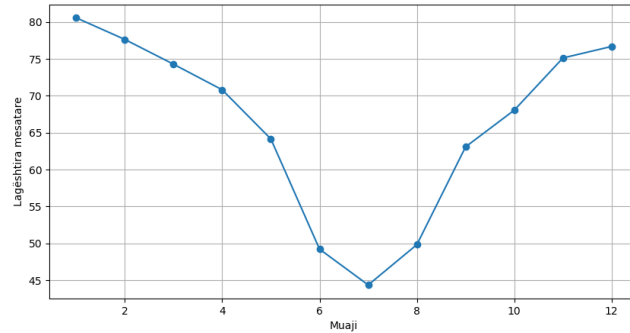


Figura 3.6 Sezonalteti mujor i lagështirës

Figura 3.7 paraqet komponentin trend të serisë kohore të lagështirës relative, i cili përfaqëson ndryshimin gradual afatgjatë të vlerave gjatë vitit. Rezultatet tregojnë një ulje progresive të lagështirës gjatë periudhës pranverë–verë dhe një rritje graduale gjatë muajve të vjeshtës dhe dimrit. Ky trend tregon se seria kohore nuk është plotësisht rastësore, por përmban strukturë temporale të rëndësishme. Identifikimi i trendit është i dobishëm për modelet e inteligjencës artificiale dhe sistemet Digital Twin, pasi mundëson kuptimin e sjelljes afatgjatë të ambientit dhe përmirëson aftësinë e sistemeve për monitorim dhe vendimmarrje automatike.

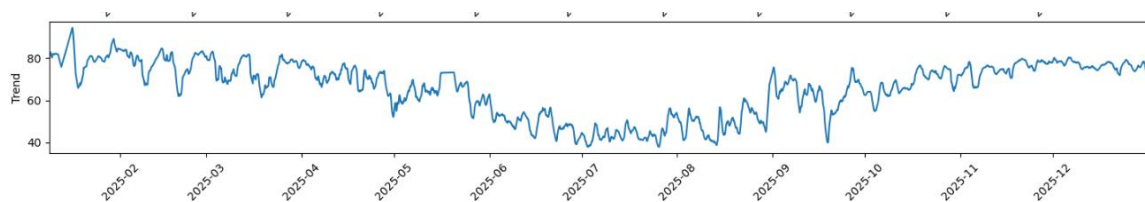


Figura 3.7 Komponenti trend i lagështirës relative gjatë vitit

3.3.6. Implikime për modelim parashikues dhe vendimmarrje automatike

Analiza e strukturës së dataset-it tregoi se të dhënat paraqesin karakteristika të përshtatshme për modelim parashikues dhe sisteme të vendimmarrjes automatike. Frekuenca e rregullt e matjeve, niveli shumë i ulët i mungesave dhe përqindja e kufizuar e anomalive tregojnë se dataset-i ruan cilësi të mirë për trajnim modelesh të inteligjencës artificiale. Analizat e autokorrelacionit dhe veçorive me vonësë kohore konfirmuan ekzistencën e varësisë kohore ndërmjet matjeve, duke justifikuar përdorimin e veçorive të vonësës dhe modeleve të serive kohore. Gjithashtu, rezultatet treguan praninë e

marrëdhënieve lineare dhe jo-lineare ndërmjet parametrave mjedisorë, çka sugjeron përdorimin e kombinimeve të modeleve statistikore, ML (Machine Learning) dhe DL (Deep Learning). Testi ADF konfirmoi stacionaritetin e serisë së lagështirës, ndërsa analiza e sezonalitetit evidentoi modele ditore dhe mujore që mund të përdoren si veçori shtesë në procesin e parashikimit. Në përgjithësi, këto karakteristika e bëjnë dataset-in të përshtatshëm për zhvillimin e modeleve parashikuese dhe integrimin e tyre në sistemin Digital Twin për monitorim dhe vendimmarrje automatike.

3.4. Analiza e të dhënave të pa strukturuara audio

Në mënyrë automatike është konvertuar formati i të gjithë audiove .mp4 dhe .m4a në format wav sepse është format i përshtatshëm për analizë. Nuk ka humbje informacioni dhe është shumë kompatibël me librari si librosa dhe torchaudio. Audio-t e regjistruara janë kategorizuar në tre grupime emocionesh, pozitive_i, ku i është numër inkrementues, negative_i si dhe neutrale_i. Secili skedar është vendosur brenda direktorive me emërtimet Pozitive_wav, Negative_wav dhe Neutrale_wav. Audio të regjistruara Pozitive janë 130, Neutrale 150 si dhe Negative 107 audio.

Tabela 3.5 Statistika përshkruese të audios

Label	Numri	Mean	Std	Min	25%	50%	75%	Max
Negative_wav	107	22.840821	4.021695	14.826688	20.010656	22.464	25.248	33.770688
Neutrale_wav	150	21.59195	5.715214	12.842688	17.130656	20.287312	25.194656	41.237313
Pozitive_wav	130	22.603799	3.754229	15.36	20.026688	22.133312	24.991016	33.726687

Tabela 3.5 paraqet shpërndarjen statistikore të kohëzgjatjes së audios (në sekonda) për tre klasat. Dataseti është i balancuar në aspektin kohor sepse mesatarja (Mean) është shumë e afërt mes këtyre klasave. Devijimi standard (std) është më i lartë për klasën neutrale, ka tregon variabilitet më të madh në kohëzgjatje krahasuar me klasat negative dhe pozitive. Vlerat minimale dhe maksimale tregojnë se klasa neutrale ka shpërndarjen më të gjerë (deri në ~41 sekonda), ndërsa klasat e tjera janë më të kufizuara. Kuartilet (25%, 50%, 75%) janë të ngjashme midis klasave, duke sugjeruar një shpërndarje relativisht të qëndrueshme dhe pa ekstreme të theksuara në pjesën më të madhe të të dhënave.

Normalizimi i audios është një hap thelbësor para analizës me modele të inteligjencës artificiale, sepse siguron që të dhënat hyrëse të jenë të standardizuara dhe të krahasueshme. Përdorimi i një frekuence të unifikuar si 16 kHz është i rekomanduar në literaturë (p.sh. wav2vec 2.0, HuBERT) sepse përmban informacion të mjaftueshëm për të folurën njerëzore duke reduktuar kompleksitetin llogaritës [86]. Konvertimi në mono shmang redundancën midis kanaleve dhe e bën modelin të fokusohet në përmbajtjen akustike kryesore [87]. Heqja e zhurmës përmirëson raportin sinjal-zhurmë (Signal to Noise Ratio-SNR), duke ndihmuar modelet të mësojnë karakteristika më të pastra të sinjalit [88], ndërsa eliminimi i heshtjes redukton të dhënat e panevojshme dhe përqendron analizën në segmentet informative të të folurit [89]. Në këtë mënyrë, normalizimi rrit performancën, stabilitetin dhe përgjithësimin e modeleve të ML (machine learning) dhe DL (deep learning) në analizën e audios.

Duke marrë në konsideratë materialin ekzistues u bë një analizë paraprake, dhe rezultatet paraqiten në tabelën e mëposhtme:

Tabela 3.6 Analiza përmbledhëse e karakteristikave akustike të dataset-it audio para normalizimit

Klasat	Nr skedarëve	Frekuenca kampionimit	Kanali	Kohëzgjatja sek			Amplituda			Raporti % i heshtjes			Vlerësimi i nivelit të zhurmës	
				min	mean	max	min	mean	max	min	mean	max	mean	max
Negative_wav	107	{16000: 107}	{1: 107}	14.826688	22.840821	33.770688	0.354431	0.639607	1	30.034239	39.388293	60.237352	0.004142	0.005004
Neutrale_wav	150	{16000: 150}	{1: 150}	12.842688	21.59195	41.237313	0.30011	0.686949	1	9.253405	39.218185	60.940032	0.004559	0.005714
Pozitive_wav	130	{16000: 130}	{1: 130}	15.36	22.603799	33.726687	0.407867	0.746946	1	10.063824	33.338482	55.393703	0.004961	0.005726

Tabela 3.6 paraqet një analizë krahasimore të tre klasave të audios (negative, neutrale dhe pozitive). Vihet re se të gjitha audiot janë të standardizuara në 16 kHz dhe mono (shumë e mirë për analizë të të folurit), çka siguron uniformitet në përpunim. Kohëzgjatja mesatare është e ngjashme midis klasave (~21–22 sekonda), duke treguar balancë në strukturën kohore të të dhënave.

Nga ana e amplitudës, klasa pozitive ka vlera mesatare më të larta, duke sugjeruar intensitet më të madh të të folurit, ndërsa klasa neutrale ka vlera më të ulëta. Raporti i heshtjes është më i lartë tek klasa negative dhe neutrale, duke treguar më shumë segmente pa zë, ndërsa audio pozitive është më “aktive”. Niveli i zhurmës është i ulët dhe relativisht i qëndrueshëm në të gjitha klasat, çka tregon cilësi të mirë të sinjalit.

Nga këto rezultate në përgjithësi duket që audiot janë në gjendje të mirë strukturore por kanë nevojë për pastrim të heshtjes para modelimit. Audio është e normalizuar deri në kufirin maksimal sepse amplituda ka vlerën max=1.0. Megjithëse mesataret e amplitudës janë në interval normal dhe nuk duket problematike mund t’i bëhet ende një normalizimi standard që të jenë njësoj. Ndërkohë heshtja është relativisht e lartë ku 40% e audios është heshtje ose shumë e ulët në amplitudë dhe duhet aplikuar heqja e zhurmës ose prerja. Niveli i zhurmës është shumë i ulët prandaj mjafton një heqje e lehtë zhurme.

Tabela 3.7 Analiza e karakteristikave akustike të dataset-it audio pas aplikimit të normalizimit

Klasat	Nr skedarëve	Frekuenca kampionimit	Kanali	Kohëzgjatja sek			Amplituda			Raporti % i heshtjes			Vlerësimi i nivelit të zhurmës	
				min	mean	max	min	mean	max	min	mean	max	mean	max
Negative_wav	107	{16000: 107}	{1: 107}	14.304	22.259239	33.12	0.999969	0.999989	1	22.631296	35.452252	51.447674	0.004177	0.005044
Neutrale_wav	150	{16000: 150}	{1: 150}	12.810688	21.220675	40.832	0.999969	0.999986	1	14.228554	38.59584	61.898216	0.004489	0.005712
Pozitive_wav	130	{16000: 130}	{1: 130}	15.328	22.319716	33.726687	0.999969	0.999989	1	12.784233	34.25002	51.586555	0.004875	0.005717

Siç vihet re (Tabela 3.7) te pjesa e nivelit të zhurmës nuk kemi ndryshime të mëdha edhe pas heqjes së zhurmës brenda intervalit të audios, por përsëri këto fragmente nuk mund të hiqen sepse janë të lidhura ngushtë edhe me emocionet e ndryshme që përcjell audio. Gjithsesi u bë një rivlerësim për shpërndarjen e heshtjes në intervalet audio për të tre klasat për ta vërtetuar këtë ndryshim të vogël.

Analiza e pauzave të brendshme tregon se çdo regjistrim përmban një numër të konsiderueshëm ndërprerjesh natyrale të të folurit, me mesatare rreth 8–10 pauza për audio në të gjitha klasat. Kohëzgjatja mesatare e këtyre pauzave varion nga rreth 0.26 deri në 0.30 sekonda, ndërsa pauzat maksimale arrijnë deri në 3.6 sekonda, veçanërisht në klasën neutrale. Në total, koha e heshtjes së brendshme përbën afërsisht 2.6–3.6 sekonda për audio, duke përfaqësuar një pjesë të rëndësishme të strukturës temporale të sinjalit. Këto rezultate konfirmojnë se pauzat e brendshme janë pjesë integrale e të folurit dhe mbartin informacion prosodik dhe emocional, ndaj ruajtja e tyre është e justifikuar në analizën e sentimentit nga audio.

3.4.1. Struktura dhe shpërndarja e të dhënave audio

Analiza e shpërndarjes së të dhënave (data distribution) për klasat e audios është një hap thelbësor në ndërtimin e modeleve të Inteligjencës Artificiale për analizën e sentimentit, sepse ndikon drejtpërdrejt në performancën dhe besueshmërinë e modelit. Nëse klasat (p.sh. pozitive, negative, neutrale) nuk janë të balancuara, modeli mund të anojë drejt klasës dominante dhe të japë rezultate të pasakta, sidomos për klasat me më pak mostra. Përmes analizës së shpërndarjes (p.sh. histogramet për numrin e mostrave dhe kohëzgjatjen e audios), mund të identifikohen probleme si imbalance, variabilitet i madh në kohëzgjatje, apo mungesë përfaqësimi e disa karakteristikave akustike. Kjo analizë ndihmon në marrjen e vendimeve të informuara për fazat pasuese, si balancimi i dataset-it (oversampling/undersampling), segmentimi i audios, ose zgjedhja e metrikave të përshtatshme të vlerësimit (p.sh. Madhësia F1 në vend të saktësisë). Në kontekstin e vendimmarrjes automatike, një shpërndarje e mirë e të dhënave siguron që modeli të jetë më i drejtë (fair), më i përgjithësueshëm dhe më i besueshëm në skenarë realë.

Për rastin konkret shpërndarja është e mirë por jo perfekte me 130 raste Pozitive, 107 raste Negative dhe 150 Neutrale. Shpërndarja e klasave është e pranueshme për trajnim të modeleve të inteligjencës artificiale, pasi nuk paraqet imbalance të theksuar. Megjithatë, ekziston një devijim i lehtë në favor të klasës neutrale, i cili duhet të merret në konsideratë gjatë fazës së trajnimit dhe vlerësimit (p.sh., duke përdorur metrika si Madhësinë F1 ose teknika balancimi nëse është e nevojshme).

Ndërkohë për segmentimin dhe modele DL (Deep Learning) përdoret shpërndarja e kohëzgjatjes ose histogrami i kohëzgjatjes. Më poshtë jepet dy histograma, njera e ndarë sipas klasave tjetra e përgjithshme.

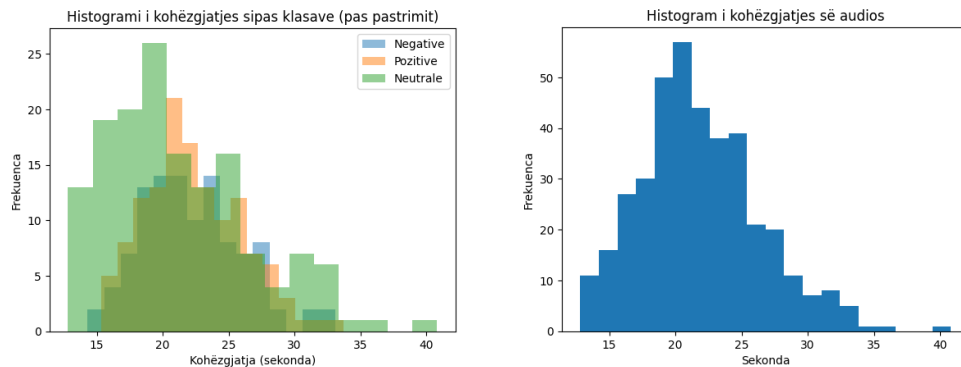


Figura 3.8 Histogram i kohëzgjatjes

Duke u nisur nga figura 3.8 shpërndarja e kohëzgjatjes së audios tregon se shumica e regjistrimeve janë të përqendruara rreth 20–25 sekonda, me një shpërndarje relativisht të qëndrueshme dhe pak raste ekstreme. Kjo strukturë e të dhënave mbështet përdorimin e segmentimit në intervale 2–4 sekonda, duke mundësuar rritjen e numrit të mostrave për trajnim dhe kapjen e variacioneve të brendshme të sinjalit audio, pa humbur kontekstin semantik. Për ta përforcuar më shumë zgjedhjen optimale të segmentimit përdoret histograma e pauzave sipas figurës 3.9:

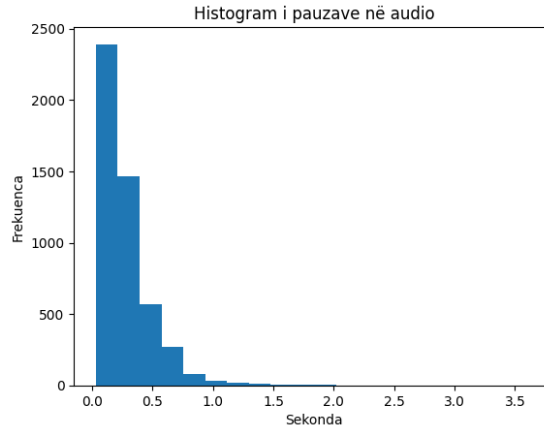


Figura 3.9 Histograma e pauzave në audio për tre klasat

Shumica e pauzave janë shumë të shkurtra (0-0.5 sekonda) dhe ka shumë pak pauza > 1 sekond, prandaj konkludohet me faktin që dataseti ka pak heshtje të panevojshme dhe strukturë natyrale të folurit, prandaj segmentimi 2-4 sekonda duket shumë i përshtatshëm. Si rezultat, dataseti konsiderohet i pastruar mirë dhe i përshtatshëm për aplikimin e teknikave të inteligjencës artificiale, pa nevojën për ndërhyrje të mëtejshme në eliminimin e heshtjes.

RMS (Root Mean Square) është një masë që tregon energjinë ose fuqinë mesatare të sinjalit audio, dmth sa “i fortë” është sinjali në një interval kohor. RMS llogaritet sipas formulës së mëposhtme:

$$RMS = \sqrt{\frac{1}{N} \sum_{i=1}^N x_i^2} \quad (3.1)$$

ku:

- x_i janë mostrat e sinjalit
- N është numri i mostrave

Interpretimi i RMS është në rastin kur vlerat e saja janë të ulëta kemi të bëjmë me heshtje ose pauzë. Një RMS mesatare nënkupton të folurën normale dhe RMS i lartë paraqet zërin e fortë ose emocionet intensive. Figura 3.10 paraqet shpërndarjen e RMS mesatar për secilën klasë përmes një boxplot-i.

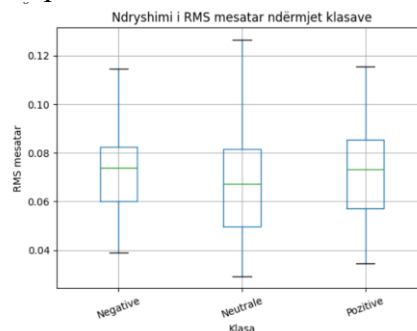


Figura 3.10 RMS mesatare ndërmjet klasave pozitive, negative dhe neutrale

Vërehet se mesataret e RMS janë të ngjashme ndërmjet klasave pozitive, negative dhe neutrale, duke treguar mungesë të dallimeve të theksuara në energjinë e sinjalit. Shpërndarja e vlerave është relativisht homogjene, çka sugjeron se të dhënat janë të normalizuara në mënyrë të përshtatshme dhe nuk paraqesin bias të rëndësishëm në lidhje me intensitetin e zërit. Impikimet për zgjedhjen e teknikave të inteligjencës artificiale Karakteristikat e identifikuara të të dhënave, përfshirë jo-linearitetin, jo-stacionaritetin dhe varësinë kohore, kanë ndikim të drejtpërdrejtë në zgjedhjen e teknikave të inteligjencës artificiale. Këto gjetje sugjerojnë se përdorimi i modeleve lineare tradicionale është i pamjaftueshëm për të kapur dinamikën komplekse të sistemit, duke justifikuar adoptimin e modeleve jo-lineare dhe adaptivë, si Random Forest dhe rrjetet nervore të thella. Kjo analizë shërben si bazë për ndërtimin e sistemeve të vendimmarrjes automatike që janë jo vetëm të sakta, por edhe të qëndrueshme dhe të besueshme në kontekste reale mjedisor.

3.4.2. Ndashmëria e veçorive

Në kontekstin e inteligjencës artificiale dhe përpunimit të sinjaleve, veçorive janë karakteristika të nxjerra nga të dhënat bruto që përfaqësojnë në mënyrë të përmbledhur informacionin më të rëndësishëm për analizë dhe modelim. Në rastin e audios, sinjali i papërpunuar është kompleks dhe me dimension të lartë, prandaj transformohet në një grup veçorish që kapin aspektet akustike dhe perceptuese të tij [90]. Sipas [85], veçoritë përbëjnë një hap kritik në pipeline-in e përpunimit të të folurit, pasi ato shërbejnë si hyrje për modelet e ML (Machine Learning) dhe ndikojnë drejtpërdrejt në performancën e klasifikimit.

Përpara përdorimit në modelet e Inteligjencës Artificiale, është e rëndësishme nëse veçoritë dallojnë klasat (emocionet pozitive, negative dhe neutrale). Ky proces njihet si Analiza e Ndashmërisë së Tipareve. Nëse një veçori ka shpërndarje të ndryshme për secilën klasë, medianet janë të dallueshme dhe ekziston një mbivendosje e vogël midis klasave, atëherë konsiderohet diskriminues dhe i dobishëm për klasifikim [91]. Pra në studim me anë të disa vizualizime janë identifikuar feature dhe jane krahasur përpara trajnimit të modelit.

Modelet e Inteligjencës Artificiale ndahen në dy kategori kryesore bazuar në mënyrën se si përdorin veçoritë.

- Modele që nxjerrin veçoritë në mënyrë manuale (ML)
- Modele që mësojnë veçoritë automatikisht (DL)

Modelet ML nuk punojnë direkt mbi sinjalin bruto dhe kërkojnë që veçoritë të nxirren paraprakisht. Raste të tilla janë LR (Linear Regression), RF (Random Forest), XGBoost dhe SVM (Support Vector Machine). Këto modele performojnë mirë kur veçoritë janë të dizenuara mirë dhe përfaqësojnë informacionin relevant.

Modelet që mësojnë veçoritë automatikisht punojnë direkt mbi të dhënat. Raste të tilla janë CNN (Convolutional Neural Network), LSTM (Long Short-Term Memory)/ GRU (Gated Recurrent Unit) dhe wav2vec 2.0, HuBERT (Hidden-Unit BERT). Sipas [80] këto modele realizojnë “mësimin e përfaqësimeve”, duke nxjerrë automatikisht veçoritë më të rëndësishme nga sinjali.

Studimet në SER (Speech Emotion Recognition) tregojnë se emocionet reflektohen në karakteristika të ndryshme akustike [84], [92].

Procesi i nxjerrjes së veçorive është realizuar duke përdorur bibliotekën *librosa*, ndërsa për të përballuar volumin relativisht të madh të të dhënave audio dhe për të optimizuar kohën e përpunimit, është aplikuar paralelizimi i proceseve përmes bibliotekës *joblib*. Në këtë mënyrë, çdo skedar audio është përpunuar në mënyrë të pavarur, duke shfrytëzuar në maksimum burimet e CPU-së (Central Processing Unit) dhe duke reduktuar ndjeshëm kohën e ekzekutimit të skriptit të nxjerrjes së veçorive. Pas nxjerrjes së veçorive, është ndërtuar një dataset i strukturuar që përfshin tre klasat e sentimentit: pozitive, negative dhe neutrale.

<i>Klasa_</i>	<i>mfcc_1</i>	<i>mfcc_2</i>	<i>mfcc_3</i>	<i>pitch</i>	<i>spectral_centroid</i>
<i>Negative</i>	282.422974	77.749336	16.335657	1223.16467	1808.03947
<i>Neutrale</i>	284.998993	104.742188	14.966753	986.472107	1499.44102
<i>Pozitive</i>	265.738037	101.210693	17.31085	1027.82239	1576.066258

Tabela 3.8 Statistika përmbledhëse e veçorive akustike për klasat emocionale

Ku:

MFCC_1 është energjia e përgjithshme e zërit

MFCC_2 është forma bazë e spektrit

MFCC_3 janë detaje të spektrit

Sa më i vogël të jetë indeksi aq më shumë informacion global përcillet kurse vlerat më të mëdha tregojnë më shumë detaje.

Pitch është lartësia e zërit

Zëri i hollë është një pitch i lartë e kundërta pitch i ulët.

Spectral Centroid është qendra e “energjisë” në spektrin e frekuencës.

Vlerat e larta të Spectral Centroidit paraqesin frekuencat e larta përndryshe frekuencat e ulëta.

Nga rezultatet e tabelës 3.8 vërehet se ekzistojnë dallime të dukshme midis klasave, veçanërisht në veçoritë si *pitch* dhe *spectral centroid*, të cilat reflektojnë ndryshime në tonalitet dhe shpërndarjen frekuenciale të sinjalit audio. Gjithashtu, koeficientët MFCC tregojnë variacione ndërmjet klasave, duke sugjeruar se ato përmbajnë informacion të rëndësishëm për karakterizimin e strukturës spektrale të të folurit. Këto dallime tregojnë se veçoritë e zgjedhura kanë potencial për të ndihmuar në dallimin e klasave dhe për rrjedhojë janë të përshtatshme për përdorim në modele të inteligjencës artificiale për analizën e sentimentit nga e folura.

Me qëllim vlerësimin e aftësisë diskriminuese të veçorive janë përdorur vizualizime grafike për të analizuar shpërndarjen e secilës sipas klasave. Një nga metodat kryesore të përdorura për këtë qëllim është boxplot, i cili lejon identifikimin e ndryshimeve në median, shpërndarje dhe praninë e outliers për secilën klasë. Përmes këtyre vizualizimeve është vlerësuar niveli i mbivendosjes midis klasave, duke ofruar një tregues të drejtpërdrejtë për ndashmërinë e veçorive. Veçoritë që shfaqin shpërndarje të dallueshme dhe mbivendosje të ulët konsiderohen më të përshtatshme për përdorim në modelet e klasifikimit.

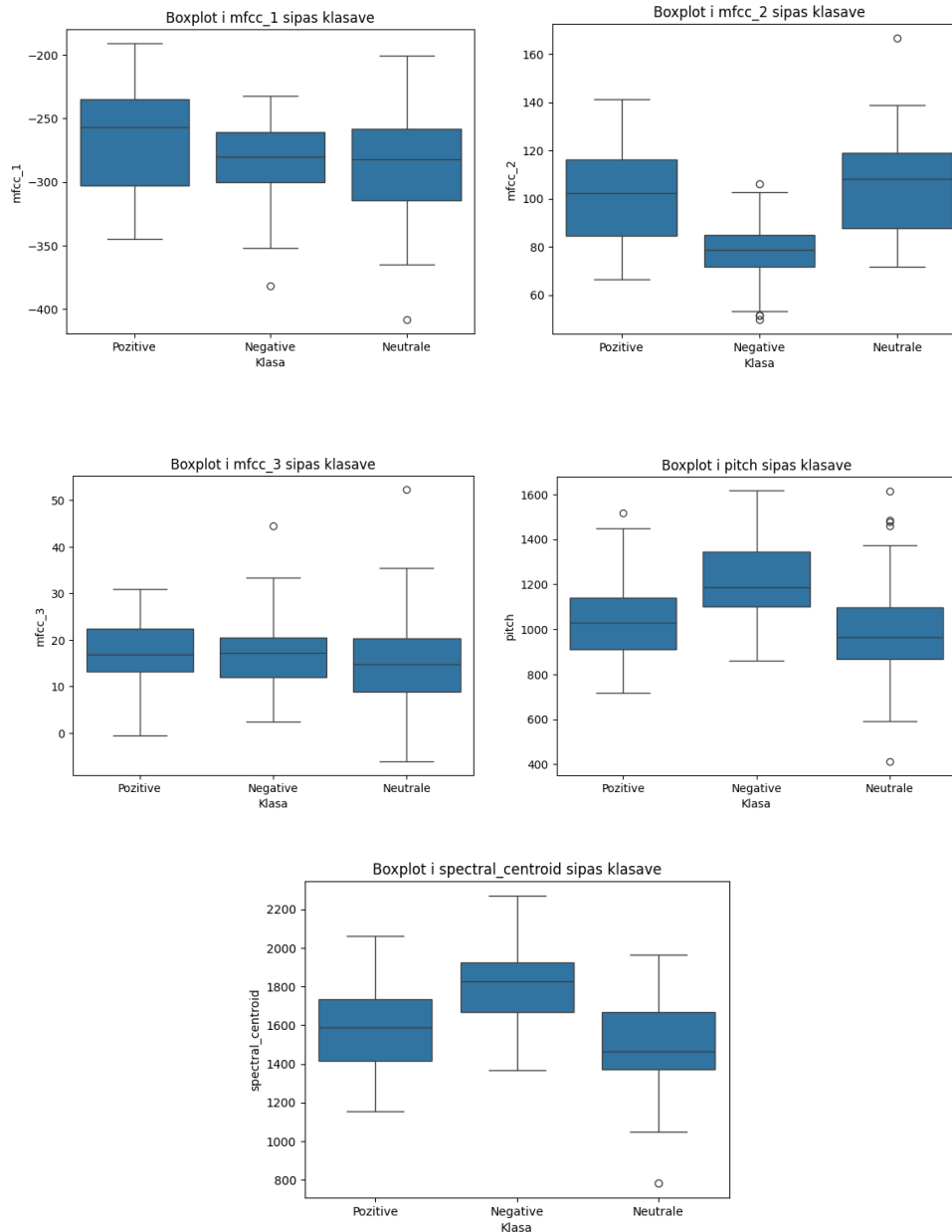


Figura 3.11 Shpërndarja e veçorive audio për vlerësimin e ndashmërisë së klasave

Analiza vizuale e shpërndarjes së veçorive e paraqitur përmes boxplot-eve të Figura 3.11, tregon se ekziston një nivel i moderuar i ndarjes midis klasave pozitive, negative dhe neutrale, megjithëse kjo ndarje nuk është uniforme për të gjitha veçoritë.

Në veçanti, mfcc_2, pitch dhe spectral centroid shfaqin dallueshmëri më të qartë midis klasave. Klasa negative paraqet tendencë për vlera më të larta në pitch dhe spectral centroid, ndërsa mfcc_2 tregon një shpërndarje më të ulët për këtë klasë krahasuar me të tjerat. Këto diferenca sugjerojnë se këto veçori kanë potencial më të lartë diskriminues dhe mund të kontribuojnë ndjeshëm në procesin e klasifikimit.

Nga ana tjetër, mfcc_1 dhe mfcc_3 paraqesin një mbivendosje më të madhe midis klasave, duke treguar se këto veçori individualisht kanë aftësi më të kufizuara për të dalluar emocionet. Megjithatë, në kombinim me veçori të tjera, ato mund të kontribuojnë në

përmirësimin e performancës së modeleve.

Në përgjithësi, rezultatet tregojnë se asnjë veçori e vetme nuk është plotësisht e mjaftueshme për ndarjen e klasave, por një kombinim i tyre krijon një përfaqësim më të pasur dhe më efektiv për modelet e ML (Machine Learning). Kjo mbështet përdorimin e një qasjeje multivariate për klasifikimin e sentimentit nga e folura.

Nxjerrja automatike e veçorive

Në dallim nga karakteristikat akustike tradicionale, si MFCC, pitch apo veçoritë spektrale, përfaqësimet e gjeneruara nga modelet e të mësuarit të thellë janë shumëdimensionale dhe nuk lidhen drejtpërdrejt me një interpretim të vetëm të kuptueshëm. Në këtë proces, sinjalet audio konvertohen në forma numerike përmes teknikave të DL, duke përdorur modele të bazuara në arkitekturën Transformer. Këto modele janë të afta të mësojnë përfaqësime të avancuar nga sinjalet audio duke kapur marrëdhënie komplekse kohore dhe spektrale. Transformer-at, të prezantuar fillimisht nga [93], mbështeten në arkitekturën self-attention, e cila mundëson identifikimin e marrëdhënieve afatgjata ndërmjet elementeve të të dhënave pa u bazuar në mekanizma rekursivë si RNN ose LSTM. Ky mekanizëm i lejon modelit të vlerësojë rëndësinë relative të secilës pjesë të hyrjes në raport me pjesët e tjera të sekuencës. Në rastin e një audio disa segmente janë më të rëndësishme dhe me anë të self-attention i jepet një peshë më e madhe këtyre pjesëve.

$$Attention(Q, K, V) = softmax((QK^T)/sqrt(d_k))V \quad (3.2)$$

Ku:

- Q (kërkesa) çfarë po kërkojmë
- K (çelësi) me çfarë po e krahasojmë
- V (vlera) informacioni real që përdoret
- softmax normalizon peshat

Çdo sinjal audio përpunohet dhe kthehet në një vektor veçorish ose karakteristikash dhe për tre klasat dallojnë 769 kolona ose veçori. Numri i rreshtave mbetet i njëjtë si në fillim. Pra, dataset-i përfaqësohet si një hapësirë me dimension të lartë [86].

Për të bërë një analizë ndërklasore aplikohen algoritmet e klasifikimit dhe reduktimit dimensional mbi një dataset të bashkuar të tre klasave.

Për të identifikuar veçoritë më të rëndësishme për dallimin e klasave, u përdor modeli Random Forest, i cili vlerëson rëndësinë e secilit atribut në procesin e klasifikimit. Rezultatet (Figura 3.12) tregojnë 10 veçoritë më të rëndësishme ku secila kontribuon në mënyrë të ndryshme në ndarjen e klasave. Disa veçori kanë peshë më të madhe (p.sh. veçoria 252) duke arritur në konkluzionin që të gjitha dimensionet kontribuojnë njësoj

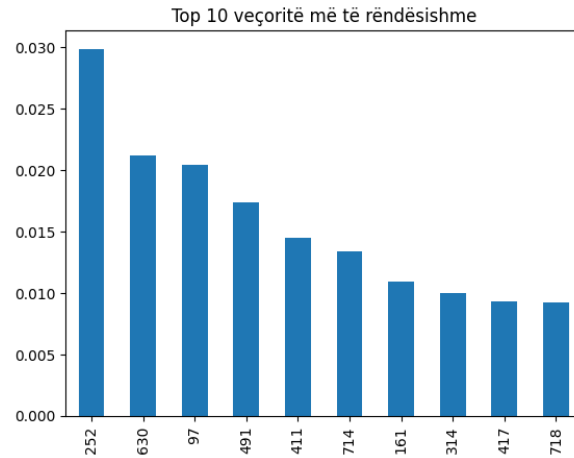


Figura 3.12 Dhjetë veçoritë me ndikim më të lartë

Për të analizuar strukturën e të dhënave në një hapësirë më të thjeshtë, e cila mund të jetë dy-dimensionale ose tre-dimensionale, u aplikua metoda PCA (Principal Component Analysis) e cila redukton dimensionet duke ruajtur variancën maksimale të të dhënave [94].

$$Z = XW \quad (3.3)$$

ku:

- X = matrica e të dhënave (rasti konkret 769 veçori përçdo audio)
- W = vektorët eigen (drejtimet kryesore ku ka më shumë informacion)
- Z = të dhënat e transformuara ose të reja

Pas aplikimit të PCA vihet re që identifikohen disa grupime por klasat nuk janë plotësisht të ndara. Ekziston një mbivendosje midis klasave pozitive dhe neutrale. Figura 3.13 paraqet këto grupime të cilat nuk janë linearisht të ndara megjithëse të dhënat kanë strukturë.

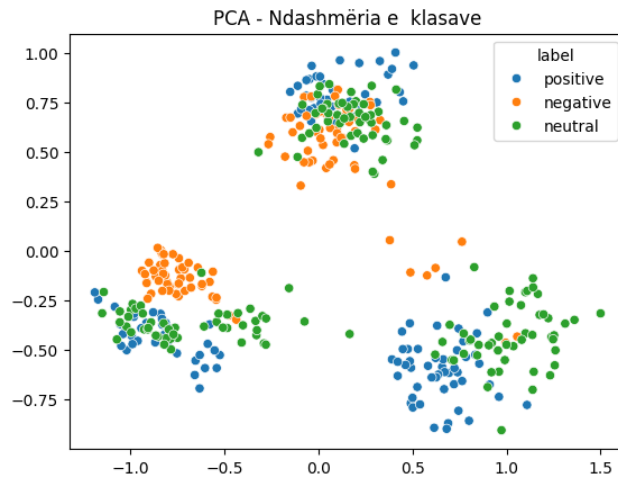


Figura 3.13 Aplikimi metodës PCA mbi të dhënat audio

Duke marrë për bazë faktin që PCA është për vizualizim, është linear dhe redukton nga 769 në 2 dimensioneve mund të humbasë informacion, dhe në rastin konkret përzierja e klasave nuk do të thotë që veçoritë e tyre nuk ndahen, thjesht në paraqitjen 2D linear nuk

është i qartë ky diskriminim. Formë tjetër për të vërtetuar që veçoritë janë të mira është F1-Score që tregon saktësinë e dallimit të klasave dhe kapjen e të gjitha rasteve [95].

$$F1 = 2 \cdot \frac{Precision \cdot Recall}{Precision + Recall} \quad (3.4)$$

Nga Tabela 3.9 rezulton se modeli arrin performancë shumë të lartë në të gjitha klasat me një saktësi prej 92%. Klasa negative paraqet performancën më të lartë, me Madhësia F1 prej 0.95 dhe precision maksimal (1.00), duke treguar se kjo klasë është më e dallueshme nga të tjerat. Ndërkohë, klasat neutrale dhe pozitive shfaqin vlera pak më të ulëta, por ende të larta (Madhësia F1 përkatësisht 0.91 dhe 0.92), duke sugjeruar një nivel të moderuar mbivendosjeje midis tyre. Mesatarja makro (macro avg) dhe ajo e ponderuar (weighted avg) konfirmojnë stabilitetin e modelit, me vlera të përgjithshme Madhësia F1 rreth 0.92–0.93, duke treguar se modeli ka aftësi të mirë përgjithësimi dhe performancë të balancuar midis klasave.

	<i>Precizioni</i>	<i>Recall (kapje e të gjitha rasteve)</i>	<i>Madhësia F1</i>	<i>Nr reali mostrave</i>
<i>Negative</i>	1	0.9	0.95	21
<i>Neutrale</i>	0.88	0.94	0.91	32
<i>Pozitive</i>	0.92	0.92	0.92	25
<i>Saktësia</i>			0.92	78
<i>macro avg</i>	0.93	0.92	0.93	78
<i>weighted avg</i>	0.93	0.92	0.92	78

Tabela 3.9 Performanca e modelit të klasifikimit për secilën klasë

Për të analizuar më mirë strukturën e të dhënave me metodë jo-lineare, u përdor t-SNE (t-distributed Stochastic Neighbor Embedding). t-SNE minimizon diferencën midis shpërndarjeve probabilitike në hapësirën origjinale dhe atë të reduktuar. Ajo kap marrëdhëniet komplekse jo-lineare [96].

Figura 3.14 paraqet shumë më qartë të gjitha grupimet, kryesisht klasa negative, e cila është e ndarë mirë nga klasat e tjera. Ka mbivendosje të vogla të klasës pozitive dhe neutrale. Pra në konkludim këto veçori janë të forta për grupimet lokale megjithëse dallimi midis disa emocioneve mbetet pak më kompleks

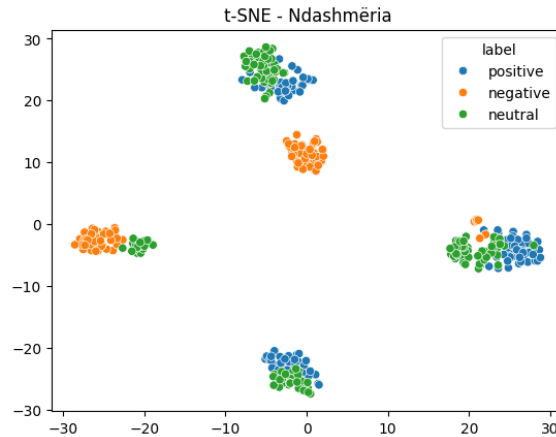


Figura 3.14 Projektioni i të dhënave në 2D me t-SNE

Transformimi i audios në vektorë numerik (embedding) me Transformer krijon një përfaqësim të pasur dhe informativ megjithëse kemi një dataset-i të strukturuar dhe jo plotësisht linearisht i ndashëm .

3.5. Lidhja me vendimmarrjen automatike

Analiza e strukturës së të dhënave përbën hallkën lidhëse midis të dhënave bruto dhe vendimeve të automatizuara. Duke kuptuar kufizimet dhe potencialin e të dhënave, është e mundur të ndërtohen mekanizma vendimmarrës që integrojnë parashikimin, vlerësimin e rrezikut dhe shpjegueshmërinë e rezultateve. Në këtë kuadër, analiza e strukturës së të dhënave nuk përbën vetëm një hap teknik, por një komponent thelbësor metodologjik për ndërtimin e sistemeve të përgjegjshme të vendimmarrjes automatike.

KAPITULLI IV

Integrimi i Shpjegueshmërisë së Inteligjencës Artificiale në ndihmë të vendimarrjes në hoteleri

Inteligjenca Artificiale është pjesë e fushave të ndryshme të jetës sonë. Mikpritja është një fushë e njohur që ndonjëherë ka nevojë për mbështetje duke adoptuar dhe integruar IA-në, siç janë chatbot-et, motorët e rekomandimeve, analizat parashikuese, përmirësimin e vendimmarrjes etj. Për [97]. IA mund të rrisë efikasitetin dhe kënaqësinë e shërbimit ndaj klientit në industrinë e mikpritjes duke shfrytëzuar rekomandime të personalizuar, sisteme automatike të mbështetjes së klientit dhe informacion të gjeneruar nga analizimi i të dhënave [98] deklaruan se IA dhe cilësia e shërbimit të punonjësve ndikojnë si në kënaqësinë dhe besnikërinë e klientit në industrinë e hoteleve, por vetëm IA nuk është e mjaftueshme për të nxitur këto rezultate. Në një studim të bazuar në të dhëna empirike nga hotelet luksoze [99] u zbulua se IA ka potencialin të rrisë kënaqësinë e klientëve në hotelet luksoze, por bashkëveprimi njerëzor është gjithashtu thelbësor për një ekuilibër. Një perspektivë tjetër e paraqitur nga [100] ishte se një zbatim i suksesshëm i IA-së varet nga integrimi i fuqisë punëtore dhe kënaqësia e punonjësve. Në sektorin e mikpritjes dhe shërbimeve, vlerësimet e klientëve luajnë një rol vendimtar në cilësinë e shërbimit dhe reputacionit të përgjithshëm. Këto komente zakonisht bëhen në platforma online si Google, ku klientët, bazuar në përvojën e tyre, mund të lënë komente që kanë një ndikim të drejtpërdrejtë dhe thelbësor në formimin e reputacionit të hoteleve [101]. Vlerësimet e klientëve online, të cilat shpesh publikohen në platforma si Google, Booking.com dhe TripAdvisor, kanë një ndikim të rëndësishëm në reputacionin, të ardhurat dhe strategjitë e mbajtjes së klientëve të një hoteli [102]. Por përpunimi dhe përgjigjja manuale ndaj këtyre vlerësimeve kërkon kohë dhe është një përpjekje e madhe, veçanërisht në mjedise me trafik të lartë. Në këtë rast studimor është përdorur inteligjenca artificiale gjeneruese duke automatizuar procesin e përgjigjes bazuar te vlerësimet e klientëve si dhe duke ruajtur një përgjigje profesionale dhe të ngjashme me atë njerëzore. Në mënyrë specifike, ne paraqesim një sistem chatbot të ndërtuar mbi modelin GPT-3.5, të integruar me mbështetje shumëgjuhëshe dhe karakteristika shpjegueshmërie të përshtatura për aplikacionet e mikpritjes. Chatbot-et tani janë bërë mjete të domosdoshme në sektorin e mikpritjes për të përmirësuar shërbimin ndaj klientit dhe për të përmirësuar operacionet e përditshme. Modelet tradicionale gjeneruese të inteligjencës artificiale, si ato të bazuara në arkitekturën transformuese, mundësojnë ndërveprim dinamik dhe personal. Në kundërshtim me intuitën, më shumë përparime në modelet gjuhësore të bazuara në transformatorë, si seria GPT e zbatuar nga OpenAI, i mundësuan chatbot-eve të trajtojnë struktura të sofistikuara gjuhësore, të mundësojnë rrjedhën natyrale të bisedës dhe të përgjigjen në kontekst. Edhe pse të pajisura me aftësi të shkëlqyera, një nga disavantazhet kryesore të modeleve të mëdha gjuhësore, LLM (Large Language Modeling) është paqartësia e tyre. Përdoruesit dhe administratorët e tyre zakonisht nuk e dinë se si dhe pse krijohen përgjigje specifike, dhe nganjëherë modelet e trajnuara keq gjenerojnë përgjigje të gabuara. Për Amann [103] kjo sjellje quhet "kuti e zezë" e cila përballlet me sfida për besimin, llogaridhënien dhe pajtueshmërinë në skenarët e shërbimit ndaj klientit. Për të kapërcyer këtë, janë zhvilluar korniza të Shpjegueshme të IA-së, XAI (eXplainable AI) për të përmirësuar dhe për t'i bërë më të shpjegueshme vendimet e sistemeve të IA-së, duke i bërë ato më transparente dhe të auditueshme për përdoruesit fundorë dhe vendimmarrësit [104]. Ndërsa modelet e

fuqishme të IA-së, si GPT, BERT ose rrjetet e thella neurale, mund të bëjnë parashikime shumë të sakta, ato shpesh funksionojnë si një "kuti e zezë" - ne marrim një rezultat, por nuk e dimë se si ose pse IA arriti në atë përfundim. XAI (eXplainable AI) ndihmon në hapjen e asaj kutie. IA e Shpjegueshme luan një rol të rëndësishëm në rritjen e besueshmërisë dhe transparencës së proceseve të vendimmarrjes të bazuara në Inteligjencën Artificiale [105]. Modelet e interpretueshme të IA-së mund të rrisin shpjegueshmërinë në vendimmarrjen komplekse pa kompromentuar performancën [106].

4.1. Chatbot-et dhe automatizimet inteligjente në turizëm

Fusha e inteligjencës artificiale brenda mikpritjes ka parë një rritje të ndjeshme, veçanërisht aplikimi i chatbot-eve për një angazhim më të madh të klientëve dhe automatizimin e shërbimeve rutinë. Shumë studime kërkimore në të kaluarën janë përqendruar në aplikimin e chatbot-eve në proceset rutinë të përditshme, siç janë ndihma gjatë rezervimeve, regjistrimi automatik, zgjidhja e ankesave ose ofrimi i informacionit të përgjithshëm në lidhje me operacionet e hotelit.

Teknologjitë automatike të vendimmarrjes që përfshijnë inteligjencën artificiale, të mësuarit automatik, automatizimin dhe inteligjencën e biznesit, prodhojnë përfitime të matshme në operacionet e mikpritjes.

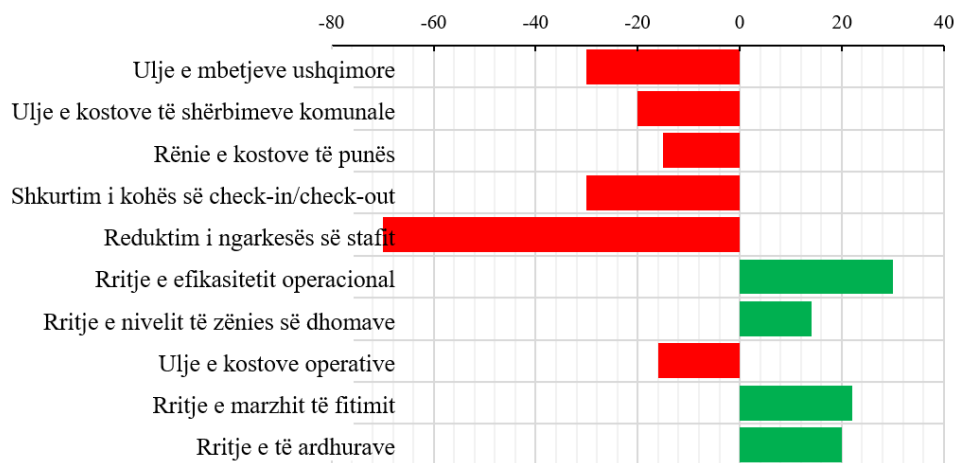


Figura 4.1 Kuantifikimi i impaktit të IA dhe Automatizimit në fushën e mikpritjes

Figura 4.1 paraqet gjithashtu metrikat operationale me variacione përqindjeje sipas disa artikujve shkencorë. Kumar dhe Ranjan [107] raportojnë kohë më të shkurtra përpunimi dhe më pak gabime në regjistrimin dhe faturimin e mysafirëve, ndërsa për Bhuiyan [108] vërehet një rritje prej 20% të të ardhurave, një rritje prej 22% të marzhit të fitimit, një ulje prej 16% të kostove operative dhe një rritje prej 14% të kapacitetit të hoteleve në formate të ndryshme hoteliere. Në mjediset e hoteleve me pesë yje, [109] vërejnë një rritje prej 30% të efikasitetit operacional, një ulje prej 70% të ngarkesës së punës së stafit dhe një shkurtim prej 30% të kohës së hyrjes/daljes; [110] regjistron një rënie prej 15% të kostove të punës, 20% kosto më të ulëta të shërbimeve dhe një rënie prej 30% të mbeturinave ushqimore. [111] dokumentojnë më tej përmirësime në saktësinë e parashikimit deri në 93% që kontribuojnë në produktivitet më të lartë. Studimet përshkruajnë përfitime cilësore në ofrimin e shërbimeve, vendimmarrje të përmirësuar dhe përvoja të përmirësuara për

mysafirët përmes teknologjive të tilla si ndërfaqet pa kontakt dhe automatizimi robotik i proceseve, por pa lënë pas dore sfidat e zbatimit, duke përfshirë kapitalin e lartë fillestar, kompleksitetet e integritetit të sistemit dhe nevojën për trajnimin e stafit. Në përgjithësi, provat mbështesin që teknologjitë e vendimmarrjes automatike japin përmirësime të prekshme në efikasitetin operacional brenda industrisë së mikpritjes. Inteligjenca Artificiale e Shpjegueshme është shfaqur për të adresuar dhe sqaruar kompleksitetin e modeleve të IA-së, veçanërisht rrjeteve të thella nervore [112]. XAI synon të zhvillojë sisteme IA të interpretueshme dhe transparente, duke rritur besimin dhe proceset e vendimmarrjes në fusha të ndryshme. Ekzistojnë teknika të përdorura për të krijuar modele të interpretueshme dhe shpjegueshmëri [107] si teknika të tilla si Shpjegime Shtuese të SHapley (SHAP) dhe Shpjegimi lokal i Modelit të Interpretueshëm -Agnostik (LIME), të cilat përdoren për interpretimin dhe shpjegimin e vendimeve të modeleve të Inteligjencës Artificiale. Ndërsa aplikacioni inteligjent bëhet më i zakonshëm, është e rëndësishme të përdoren mjete që i ndihmojnë njerëzit të kuptojnë qartë se si merren vendimet. Krahasimi i realizuar mes SHAP dhe LIME për Roshinta [113] thekson faktin se SHAP ofron njohuri më të forta dhe globale sesa LIME për IA të shpjegueshme, pavarësisht kostos më të lartë llogaritëse.

4.1.1. Chatbot-et në mikëpritje

Chatbot është një sistem aplikativ që kryen biseda në një mënyrë interaktive. Një studim nga Cameron R. Jones dhe Benjamin K. Bergen raportoi se GPT-4.5 u identifikua si njeri në 73% të rasteve, ndërsa LLaMa 3.1 në 56%. Autorët e konsiderojnë këtë si prova e parë empirike që një sistem IA kalon një standard Test Turingu me tre palë.[114] gjithsesi te [115] shprehet se një LLM mund të kalojë disa versione të Testit të Turing, por dështon në versione më robuste, me kritere më të forta, kohë, më të gjatë bisede, vlerësues me eksperiencë, dhe mjedise më të strukturuar. Chatbot-et përdorin kuptimin dhe përpunimin e gjuhës natyrore për të ndihmuar përdoruesit me detyra të ndryshme [116]. Ata bashkëveprojnë me përdoruesit duke përdorur gjuhë natyrore ose tekst, duke përdorur algoritme të IA-së për të gjeneruar përgjigje [117]. Chatbot-et kanë evoluar ndjeshëm që nga fillimi i tyre, nga sistemet e hershme të bazuara në rregulla si ELIZA në vitet 1960 deri te agjentët e sotëm të përparuar të bisedave të mundësuar nga IA si ChatGPT [118], [119]. Zhvillimi i chatbot-eve është nxitur nga përparimet në përpunimin e gjuhës natyrore, të mësuarit automatik dhe inteligjencës artificiale [120]. Zhvillimet kryesore në këtë fushë përfshijnë konceptimin e testit Turing dhe realizimin e projekteve me ndikim të madh, si CALO (Cognitive Assistant that Learns and Organizes), i cili konsiderohet një ndër projektet më të rëndësishme në historinë e asistentëve inteligjentë dhe chatbot-eve modernë, duke kontribuar në evolucionin e modeleve të bazuara në arkitekturën Transformer. Chatbot-et kanë gjetur aplikime në sektorë të ndryshëm, duke përfshirë kujdesin shëndetësor, arsimin, argëtimin dhe tregtinë [118]. Ndërsa chatbot-et bëhen më të përhapura në situata të përditshme, siç janë ndërveprimet e tregtisë elektronike, zhvillimi i tyre ngre shqetësime etike në lidhje me paragjykimin, transparencën dhe privatësinë. Zhvillimet e ardhshme në teknologjinë e chatbot-eve pritet të përqendrohen në ndërveprimet e personalizuar për të rritur angazhimin dhe kënaqësinë e përdoruesit [119]. Procesi fillon me një pyetje të përdoruesit si, "A ka dhoma të lira fundjavën e ardhshme?" e cila përkufizohet si një ndërveprim midis klientit dhe sistemit. Hapi tjetër është Chatbot-

i Hibrid. Quhet Chatbot Hibrid sepse është një kombinim i një komponenti të Bazuar në Rregulla, i cili kërkon në bazën e të dhënave për të marrë informacionin e kërkuar nga përdoruesi dhe komponentit të IA-së, i cili gjeneron gjuhë natyrore në mënyrë që përdoruesi të marrë një përgjigje sa më të kuptueshme të jetë e mundur. Hapi tjetër është procesi i integritimit të komponentit XAI (eXplainable Artificial Intelligence), konkretisht SHAP (SHapley Additive exPlanations). Komponenti shpjegon pse një përgjigje e caktuar është dhënë nga chatbot-i. Kjo rezulton në një rritje të besueshmërisë në sytë e përdoruesit, si dhe qartësi të përgjigjes në aplikimet ku klienti merr vendime të rëndësishme. Së fundmi, është hapi i fundit i përcaktuar si “Përgjigje”, i cili ofron një përgjigje të personalizuar, "Kemi dhoma të disponueshme të premten dhe të shtunën. Dëshironi një dhomë standarde apo luksoze?”.

Shumë sektorë po eksplorojnë mënyra për të zbuluar dhe për të shfrytëzuar fuqinë e ChatGPT dhe mjeteve të ngjashme në bizneset e tyre, një prej të cilave është mikpritja. Studimet e fundit tregojnë se, ndërsa industria e mikpritjes dhe turizmit ka qenë gjithmonë prapa në drejtim të miratimit të teknologjive të reja për shkak të karakterit të orientuar drejt shërbimit të operacioneve të saj, ekziston një tendencë në rritje drejt dixhitalizimit [121]. Mjete unike si ChatGPT, kanë mahnitur me aftësinë e tyre për të përmirësuar përvojën e klientit dhe produktivitetin e punonjësve përmes përgjigjeve të shpejta dhe të sakta ndaj pyetjeve të përgjithshme [122]. Këto sisteme i ndihmojnë vizitorët të marrin vendime më të informuara, të tilla si rregullimet e udhëtimit, planet e udhëtimit dhe çmimet, duke përmirësuar kështu përvojën e një vizitori. Kompanitë po i integrojnë këto chatbot tani në shërbimet e pritjes dhe të portierit, rezervimet, pyetjet e mysafirëve dhe sistemet e reagimeve. Një shembull është përdorimi i ChatGPT, i cili është në gjendje t'u përgjigjet pyetjeve të shpeshta të klientëve me një shkallë të lartë saktësie dhe rrjedhshmërie, duke përmirësuar kështu përvojën e përdoruesit dhe duke liruar stafin njerëzor për probleme më komplekse [122]. Mjetet i ndihmojnë vizitorët të marrin vendime të informuara për planifikimin e udhëtimit, përzgjedhjen e hotelit dhe koston duke gjeneruar rekomandime të personalizuara sipas kërkesës.

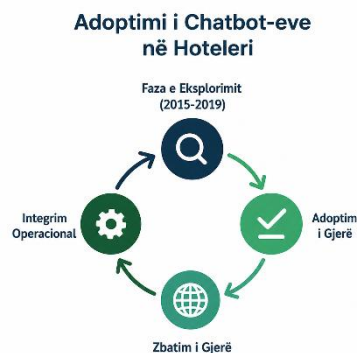


Figura 4.2 Cikli i jetës së chatbot-it

Figura 4.2 tregon përdorimin e chatbot-eve në industrinë e mikpritjes ose ndryshe hotelerisë. Zhvillimi i chatbot-eve në industrinë e mikpritjes mund të ndahet në tre faza kryesore. Gjatë fazës së parë (2015–2019), chatbot-et u përdorën në një mënyrë të kufizuar, kryesisht për pyetje të shpeshta dhe rezervime të thjeshta. Teknologjitë e përdorura përfshinin sisteme të bazuara në rregulla dhe aktivizues fjalësh kyçe, por mungesa e

fleksibilitetit dhe ndërveprimit natyror e bëri përvojën më pak tërheqëse për përdoruesit. Faza e dytë (2020–2022) pa një përshpejtim dixhital për shkak të pandemisë COVID-19, duke nxitur kërkesën për shërbime pa kontakt. Chatbot-et u përdorën për të automatizuar procese të tilla si regjistrimi/dalja, si dhe për të ofruar informacion mbi masat e sigurisë. Teknologjitë evoluuan për të përfshirë përpunimin e gjuhës natyrore (NLP) dhe mësimin automatik për analiza më të sakta të pyetjeve. Marka të njohura si Marriott, Hilton dhe Accor filluan të përdorin chatbot-e në aplikacionet e tyre mobile [123]. Faza e tretë (2023–deri më sot) karakterizohet nga integrimi i inteligjencës artificiale gjeneruese, veçanërisht përdorimi i modeleve të mëdha gjuhësore (LLM-Large Language Modeling) si GPT-5.x të OpenAI, Claude 4 / Claude Opus të Anthropic, Gemini 2.x / Ultra të Google etj. Këto modele mundësojnë përgjigje natyrale, shumëgjuhëshe dhe të personalizuar. Një aspekt i rëndësishëm është përfshirja e Inteligjencës Artificiale të Shpjegueshme (XAI), e cila justifikon vendimet e marra nga chatbot-et, duke rritur transparencën dhe besimin. Kompani të tilla si Holiday Inn (me chatbot-in Bebot), Booking.com dhe Airbnb tashmë i kanë integruar këto sisteme [123].

4.1.2. Chatbot-et gjeneruese dhe modelet e mëdha gjuhësore (LLMs)

Me kalimin e viteve, chatbot-et gjenerues kanë evoluuar nga chatbot-et e hershëm që mbështeteshin në përgjigje të paracaktuara dhe në mbështetjen e rregullave për fjalët kyçe në chatbot-e më të rinj që përdorin teknika të të mësuarit të thelluar për të prodhuar përgjigje dinamike, të ngjashme me njerëzit, të cilat janë të vetëdijshme për kontekstin dhe gjuhësisht koherente. Këto modele trajnohen në një korpus masiv teksti dhe përdorin shpërndarje probabiliteti për të parashikuar fjalën tjetër në një sekuençë, duke i mundësuar atyre të ndërtojnë përgjigje të reja në kohë reale. Chatbot-et përpunojnë të dhëna për të ofruar përgjigje për kërkesa të të gjitha llojeve duke përdorur inteligjencën artificiale, rregulla të automatizuara, përpunimin e gjuhës natyrore (NLP) dhe të mësuarit e makinës. Chatbot-et dhe modelet multimodale të bazuara në Inteligjencë Artificiale u ofrojnë kompanive të udhëtimit dhe turizmit mundësinë për të përmirësuar cilësinë e shërbimeve dhe ndërveprimit me klientët përmes automatizimit dhe optimizimit të proceseve të ndryshme. Këto teknologji mund të përdoren për mbështetjen e klientëve, gjenerimin e përmbajtjes, analizën e trendeve të tregut, zhvillimin e strategjive konkurruese, krijimin e paketave promocionale dhe trajnimin e stafit. Përmes përgjigjeve të shpejta, reduktimit të kohës së pritjes, mbështetjes shumëgjuhëshe dhe uljes së kostove operative, chatbot-et kontribuojnë në rritjen e efikasitetit dhe përmirësimin e përvojës së përdoruesit në sektorin e turizmit dhe udhëtimeve. Sipas [124] Modelet e Mëdha Multimodale (LMM) janë një hap i madh që IA të bëhet më e saktë, më krijuese, më kontekstuale dhe më e aftë për të imituar ndërveprimin njerëzor. IA gjeneruese përfaqëson një kategori të gjerë teknologjish të Inteligjencës Artificiale që përdoren për krijimin e përmbajtjes së ndryshme. Modelet e mëdha gjuhësore (LLM) kanë aftësinë të përpunojnë dhe gjenerojnë të dhëna tekstuale, ndërsa modelet multimodale mund të trajtojnë dhe prodhojnë forma të ndryshme të dhënash, si tekst, imazh, audio apo video. Chatbot-et e bazuara në IA dhe modelet multimodale ofrojnë avantazhe të shumta, përfshirë shpejtësinë, saktësinë, kreativitetin, efikasitetin dhe fleksibilitetin, duke i bërë të zbatueshme në fusha të ndryshme, si krijimi i përmbajtjes, asistenca inteligjente, mbështetja ndaj klientëve dhe optimizimi i proceseve të biznesit. [125].

Disa chatbot gjenerues të teknologjisë së fundit sot janë ndërtuar mbi modele të mëdha gjuhësore (LLM) që mundësojnë ndërveprime shumë kontekstuale dhe të ngjashme me njerëzit. Shembuj të shquar përfshijnë ChatGPT nga OpenAI, i mundësuar nga GPT-3.5 dhe GPT-4, dhe më vonë nga modele më të avancuara të gjeneratës GPT-5, i përdorur gjerësisht në fusha të tilla si arsimi, kujdesi shëndetësor dhe mikpritja. Google Board, bazuar në modelin PaLM 2 (Modeli i Gjuhës së Shtigjeve), shkëlqen në detyrat e rikthimit të informacionit dhe arsyetimit. Claude, i zhvilluar nga Anthropic, përdor IA kushtetuese për të gjeneruar përgjigje të harmonizuara, të sigurt dhe koherente. Modelet LLaMA (Meta AI i Modelit të Gjuhës së Madhe) me burim të hapur janë vënë në dispozicion nga Meta që shërbejnë si bazë për sistemet chatbot shumëgjuhëshe dhe të trajnueshme [126]. Përveç kësaj, IA gjeneruese zbatohet në softuerët e mbështetjes dhe produktivitetit të ndërmarrjeve përmes Amazon Q dhe Microsoft Copilot, duke treguar se si LLM-të do të nxisin vendimmarrjen e biznesit dhe do të përmirësojnë detyrat që kërkojnë shumë ndërveprim. Këto inovacione në inteligjencën artificiale gjeneruese janë veçanërisht të rëndësishme për mikpritjen, ku kërkesa për pritjet e mysafirëve, për shërbim dixhital në kohë reale, të personalizuar dhe të besueshëm, janë vazhdimisht në rritje.

4.1.3. IA e Shpjegueshme (XAI) dhe Sistemet e bashkëbisedimit

Në fushat e mikpritjes dhe shërbimit ndaj klientit, besimi në sistemet e IA-së varet jo vetëm nga rrjedhshmëria e përgjigjeve, por edhe nga aftësia e përdoruesve për të kuptuar pse është marrë një vendim ose rekomandim i caktuar. Kështu, nevoja për transparencë dhe interpretueshmëri në proceset e tyre të vendimmarrjes bëhet gjithnjë e më kritike. Për ta mundësuar këtë, Inteligjenca Artificiale e Shpjegueshme (XAI) na vjen në ndihmë. Një nga objektivat kryesore të Inteligjencës Artificiale të Shpjegueshme (XAI) është rritja e besimit ndaj teknologjisë përmes ofrimit të transparencës dhe kuptueshmërisë në procesin e vendimmarrjes. XAI synon t'i mundësojë përdoruesit të marrë informacion dhe shpjegime të drejtpërdrejta nga sistemi inteligjent lidhur me rezultatet e gjeneruara. Në këtë kontekst, XAI përfshin metoda dhe teknika që e bëjnë funksionimin e modeleve komplekse të inteligjencës artificiale më të interpretueshëm dhe më të kuptueshëm për njerëzit. [127]. Teknologjitë tradicionale të IA-së, siç është të mësuarit e thellë (p.sh., GPT), janë zakonisht "kutitë e zeza" në kuptimin që ajo që ndodh brenda kompanisë nuk është e kuptueshme. Sfidat e moskuptimit të tyre është një problem për firma të tilla si mikpritja, kënaqësia e klientit dhe zgjedhjet operacionale të cilave duhet të jenë transparente, të auditueshme dhe në përputhje me politikat e firmës [70]. Chatbot-et XAI do të jenë në gjendje të justifikojnë rezultatet, p.sh., pse një klient nuk mund të zhvendosej në një dhomë ose pse u sugjerua një shërbim i caktuar. Në fushat që kërkojnë saktësi dhe transparencë të lartë, siç janë mikpritja, kujdesi shëndetësor ose financa, shpjegueshmëria e sistemeve të bisedës është një element parësor, qoftë për debugging dhe auditimin e sjelljes së modelit, apo ndërtimin e besimit me përdoruesin. Në mënyrë specifike, një chatbot gjenerues që ndihmon klientët e hotelit duhet të mbulojë një gamë funksionalitetesh, duke filluar nga informimi i përdoruesit nëse një përmirësim dhome është i disponueshëm, deri te shpjegimi i arsyes pas çdo përgjigjeje, siç është statusi i nivelit të besnikërisë ose disponueshmëria e dhomës. Një nivel i tillë transparence i ndihmon përdoruesit ta perceptojnë sistemin si më të drejtë dhe më të ditur. Sipas [128] efektiviteti i inteligjencës artificiale të shpjegueshme në sistemet bisedore varet nga saktësia e shpjegimit dhe përputhshmëria që shpjegimi ka

me modelin mendor të përdoruesit, veçanërisht në rastet kur përdoruesi nuk është i njohur me inteligjencën artificiale. Studimi i këtyre elementëve thekson se paraqitja e shpjegimeve duhet të jetë bisedore dhe e ndjeshme ndaj kontekstit për të arritur një nivel besimi për kënaqësinë e përdoruesit. Në aplikacionet e mikpritjes, një numër i lartë përdoruesish mund të mos jenë të njohur me anën teknike, dhe si pasojë shpjegimet bisedore ndihmojnë në sqarimin pse një chatbot rekomandon ose refuzon shërbime, të tilla si përmirësimet e dhomave ose ofertat e personalizuara. Në artikullin [128] theksohet se shpjegimet interaktive nxisin një kuptim më të madh, duke u lejuar përdoruesve të bëjnë pyetje ose të kërkojnë sqarime të mëtejshme, duke i bërë sistemet e inteligjencës artificiale të perceptohen si më transparente dhe të përgjegjshme. Kjo formë bisedore e IA-së është veçanërisht kritike në fusha të shërbimit siç është mikpritja, ku përvoja e përdoruesit dhe drejtësia e perceptuar janë parësore për një përvojë të kënaqshme të përdoruesit. Duke u bazuar në këtë, [129] u krahasua ndikimi i ndërfaqeve bisedore XAI me atë të paneleve tradicionale XAI. Bazuar në raportin e tyre, ndërsa ndërfaqet bisedore rrisin paksa mirëkuptimin dhe besimin e përdoruesve, ato gjithashtu shkaktajnë mbështetje të tepërt, veçanërisht kur mundësohen nga agjentë të bazuar në LLM. Studimi thekson se këto lloje ndërfaqesh i bëjnë njerëzit të mendojnë se kanë një përshtypje të gabuar për aftësinë e IA-së, atë që autorët e quajnë "iluzioni i thellësisë shpjeguese". Këto gjetje nxjerrin në pah rëndësinë e hartimit me kujdes të sistemeve bisedore XAI në mënyrë që ndërveprimi dhe të menduarit kritik të jenë të balancuara mirë në mënyrë që përdoruesit të mos harrojnë kufizimet e IA-së edhe kur bashkëveprojnë me ndërfaqe në dukje natyrale të ngjashme me dialogun.

4.2. Chatbot Hibrid me XAI në Hoteleri

4.2.1. Pasqyra e përgjithshme e sistemit

Sistemi i propozuar, siç përshkruhet në Figura 4.3, është një chatbot hibrid, që nuk bazohet vetëm në rregulla ose informacione të ofruara nga baza e të dhënave, por përfshin gjithashtu një chatbot të shpjgueshëm për përgjigje automatike të rishikimeve në mikpritje dhe vendime të thjeshta për shërbimin ndaj mysafirëve, duke përfshirë "Pse-në" pas secilës përgjigje. Në frontend përdoret React dhe backend Node.js. Për parashikim dhe shpjegim konsiderohet një motor vendimmarrjeje i lehtë ML (RandomForest+SHAP) i orkestruar me një shtresë gjeneruese NLG (GPT-3.5).

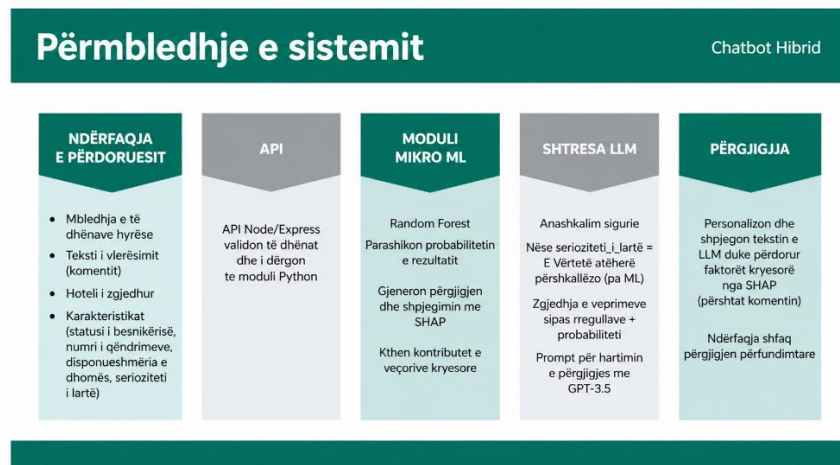


Figura 4.3 Përmbledhje e Sistemit Hibrid të Chatbot-it

4.2.2. Ndarja Funksionale në Shtresa e Sistemit

Arkitektura e sistemit të propozuar është e ndarë në tre shtresa.

1. Shtresa e prezantimit përfshin pjesën e përparme, ose ndërfaqen grafike, si një pjesë e rëndësishme që mbështet komunikimin midis klientit dhe chatbot-it. Sistemi i softuerit funksionon në shfletues. Brenda shtresës së prezantimit u përfshi React si një bibliotekë JavaScript, që rrit shpejtësinë e komponentëve të ndërfaqes grafike. Për krijimin e një front-end modern dhe të shpejtë, u përdor Vite, i cili përfshin module ES native dhe zëvendësim të moduleve të nxehta (HMR). Në çdo ndryshim të ndërfaqes grafike, Vite e rifreskon shfletuesin menjëherë. I fundit është Tailwind, i cili e bën stilimin më të shpejtë dhe të qëndrueshëm dhe është një kornizë CSS (Cascading Style Sheets) e bazuar në shërbimet e para. Në ndërfaqen grafike informacioni i mbledhur është teksti i vlerësimit, hoteli i zgjedhur dhe të dhënat informative rreth përdoruesit (p.sh. statusi i besnikërisë, numri i qëndrimeve, dhomat e disponueshme dhe ashpërsia e lartë).

2. Shtresa logjike do të thotë Back-end. Është truri i sistemit të aplikacionit ku përpunohen kërkesat e marra nga front-end. Këtu zbatohen gjithashtu rregullat e biznesit dhe komunikimi me bazën e të dhënave ose modelin e IA-së. Node.js përdoret për ndërtimin e aplikacioneve nga serveri. Ndryshe nga front-end-i që funksionon në shfletues, Node.js funksionon në server. Për të thjeshtuar krijimin e REST API-ve (Representational State Transfer Application Programming Interface), përdoret express.js si një kornizë web për Node.js dhe mjafton vetëm të përcaktohen rrugët (pikat fundore) në vend që të shkruhet kodi i papërpunuar i serverit. Node.js dhe Express veprojnë si një shtresë e ndërmjetme midis front-end-it (klientit) dhe skriptit ML. Rest API specifikon rrugën dhe pranon kërkesa (nga klienti) dhe dërgon përgjigje. Përmes orkestrimit, shërbime të ndryshme të këtij sistemi softuerik koordinohen.

3. Mikromoduli ML (Python) paraqitet si një komponent i zgjeruar i backend-it, i specializuar për parashikim dhe shpjegueshmëri. Ai përfshin modelin parashikues dhe modelin shpjegues. Ky është një modul i vogël por shumë i specializuar që përdoret nga serveri për të bërë parashikime të shpejta lokale. Pjesët kryesore të këtij moduli përfshijnë Modelin e Pyllit të Rastësishëm (Random Forest), një algoritëm të fuqishëm që kombinon shumë pemë vendimesh. Rezultati i këtij modeli në rastin tonë është probabiliteti ose klasifikimi i klientit. Ndërkohë, brenda këtij moduli, përcaktohet edhe metoda shpjeguese

e parashikimit të bërë, SHAP (Shapley Additive exPlanations). Shpjegimi përfshin prezantimin e peshës ose kontributit të secilës veçori në vendimmarrje. Për shembull, mund të thotë: "Statusi i besnikërisë kontribuoi 40% në vendimin e përmirësimit", "Teksti i rishikimit kontribuoi 30%" dhe "Disponueshmëria e dhomës kontribuoi 20%". Serveri Node.js e thërret këtë modul përmes një API të brendshëm. Ky modul është i lehtë dhe shërben vetëm për konkluzion (jo për trajnimin e modelit). Kjo qasje quhet mikro-modul sepse është një shërbim i vogël dhe i pavarur që kryen parashikime dhe shpjegime. Skedari `ml_responce_generator.py` është truri lokal i chatbot-it. Përmes tij, parashikimi, vendimmarrja dhe shpjegimi kryhen me ndihmën e SHAP (SHapley Additive exPlanation).

4.2.3. Moduli shpjegueshmërisë (Integrimi i XAI)

IA e shpjgueshme ka të bëjë me kthimin e konceptit të inteligjencës artificiale të zbehur si "kutie misterioze". Ndonjëherë një përgjigje nuk është e mjaftueshme, por ne duam të dimë pse është dhënë kjo përgjigje, cilët faktorë e kanë ndikuar atë dhe sa e besueshme është. Kjo është arsyeja pse inteligjenca artificiale funksionon më mirë duke përfshirë shpjgueshmërinë, veçanërisht për të nxjerrë në pah paragjykimet ose gabimet e fshehura. DARPA (Agjencia e Projekteve të Kërkimit të Avancuar të Mbrojtjes) e përmbledh mirë konceptin: XAI duhet ta lejojë IA-në të shpjegojë arsyetimin e saj, të tregojë pikat e forta dhe të dobëta të saj dhe t'i ndihmojë njerëzit ta kuptojnë dhe t'i besojnë asaj [130]. Cilat janë arsyet e përdorimit të XAI-së:

- Rritja e besimit ndërmjet njerëzve dhe inteligjencës artificiale
- Sigurimi që inteligjenca artificiale respekton ligjet dhe rregullat
- Identifikimi dhe reduktimi i paragjykimeve ose rezultateve të padrejta
- Justifikimi i vendimmarrjes
- Kontrolli dhe vlerësimi i funksionimit të modelit

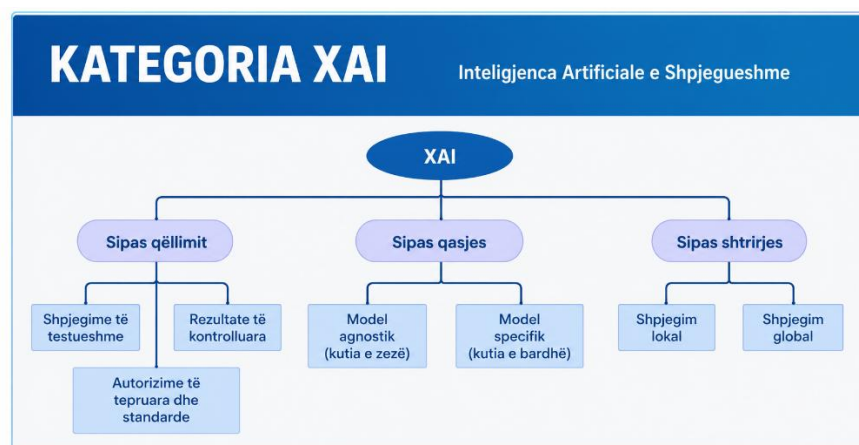


Figura 4.4 Kategorizimi i inteligjencës artificiale të shpjgueshme

Sipas [130], XAI u kategorizua Sipas Qëllimit, Sipas Fushës së Veprimit dhe Sipas Qasjes (Figura 4.4). Duke marrë parasysh kategorinë Sipas Qëllimit, interesi është t'i përgjigjemi pyetjes "pse" duam një shpjegim nga IA, prandaj rezultatet e saj janë provë, në mënyrë që gjithçka të jetë e qartë, e testueshme dhe e gjurmueshme. Pastaj janë Validimet, një nënkategori e Sipas Qëllimit, e cila verifikon nëse vendimi i marrë nga IA ka kuptim,

dhe së fundmi, janë Autorizimet që përfshijnë shpjegime që bazohen në rregulla, standarde që lidhen me përdorimin e sigurt dhe të përgjegjshëm të IA-së. Në kategorinë e dytë, "Sipas Qasjes", fokusi është në atë se sa shumë mund të shohim brenda modelit të IA-së. Rastet grupohen në modelin agnostik, ku vendimi shpjegohet pa pasur nevojë të kuptohet funksionaliteti i modelit, dhe duke qenë të kënaqur me të dhënat hyrëse dhe dalëse. Ndërkohë, modeli specifik, i njohur edhe si "Kutia e Bardhë" (White Box), përfshin raste ku kërkohet qasje në brendësi të modelit, parametrat dhe vendimet logjike. Në kategorinë e tretë, "Sipas Fushës së Veprimit", fokusi është në atë se sa i gjerë është shpjegimi. Ndarja e kësaj kategorie përfshin shpjegimin lokal, i cili përshkruan pse IA mori një vendim të caktuar për një rast, dhe shpjegimin global, i cili përshkruan logjikën ose sjelljen e përgjithshme të modelit të IA-së për të gjitha rastet.

Integrimi i IA-së së Shpjegueshme realizohet duke shtuar një modul post-hoc në chatbot-in ekzistues. Ky modul funksionon paralelisht me motorin kryesor të dialogut të chatbot-it. Motori i dialogut kupton se çfarë kërkon përdoruesi dhe jep përgjigjen ose vendimin. Në vend të kësaj, moduli XAI (post-hoc) analizon vendimin dhe gjeneron një shpjegim. Shpjegimet e ofruara kategorizohen si globale dhe lokale. Siç u përmend nga LIME (Local Interpretable Model-agnostic Explanation), shpjegimet lokale ofrojnë interpretim për secilën përgjigje individuale duke modeluar sjelljen lokale të modelit. Shpjegimet globale tentohen nga SHAP (Shapley Additive exPlanation) dhe ndihmojnë në identifikimin se si karakteristikat e ndryshme ndikojnë në situatën e përgjithshme [131]. Kjo qasje e bën chatbot-in më transparent dhe ndihmon në ndërtimin e besimit të përdoruesit, pasi ata mund të kuptojnë "pse" është paraqitur një përgjigje e caktuar, p.sh. "Sugjerimi ishte bazuar në zgjedhjet e mëparshme të përdoruesve të ngjashëm" [132]. Për më tepër, metoda dialoguese XAI, e përshkruar në literaturë si një qasje "me shumë kthesa", mundëson eksplorim interaktiv ku përdoruesit mund të bëjnë pyetje shtesë dhe të thellojnë shpjegimet në një stil bisedor [133].

4.2.4. Implementimi praktik dhe ndërfaqja

Sistemi i propozuar integron një modul të Inteligjencës Artificiale të Shpjegueshme për të "përmirësuar shpjegueshmërinë, interpretueshmërinë dhe besimin midis përdoruesve në procesin automatik të përgjigjes ndaj vlerësimeve". Në vendosjen e saj aktuale, shtresa e shpjegueshmërisë përdor një strategji arsyetimi të bazuar në rregulla për të futur shpjegime në çdo përgjigje të gjeneruar. Për shembull, nëse sistemi zbulon një ndjenjë negative në një vlerësim, siç është "Jo miqësore dhe jo profesionale në sportelin e kontrollit", sistemi gjeneron automatikisht një përgjigje kërkimfaljeje me një përshkrim se pse është bërë kërkimfalja, duke e lidhur me ndjenjën negative. Në mënyrë të ngjashme, nëse sistemi zbulon një ndjenjë pozitive në një koment, siç është "Një hotel i mrekullueshëm, suitat dhe pamjet janë të jashtëzakonshme", sistemi gjeneron "Faleminderit shumë për fjalët tuaja të mira! Jemi të emocionuar që ju qëndruat në hotelin tonë dhe se suitat, pamjet dhe zona e spa-së tona i kanë plotësuar pritjet tuaja", me arsyen se komenti pozitiv shkaktoi

mirënjohjen e gjeneruar.

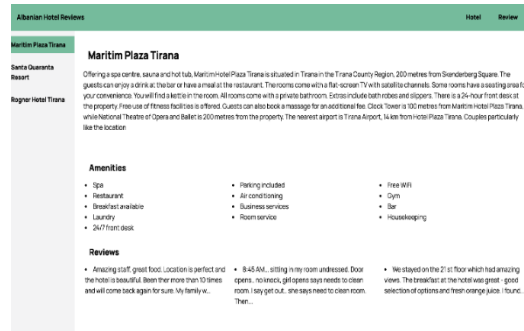


Figura 4.5 Ndërfaqja e përdoruesit

Këto shpjegime të bazuara në rregulla, të sakta, u lejojnë përdoruesve fundorë dhe menaxherëve të hoteleve të kuptojnë pse është prodhuar një përgjigje specifike, dhe kjo është thelbësore në mjediset hoteliere ku besimi dhe kënaqësia e klientit janë kritike. Sistemi praktik u zhvillua për të mbledhur dhe përpunuar vlerësimet e hoteleve. Figura 4.6 tregon ndërfaqen e parë për përdoruesin.

Siç vihet re, ekziston një listë e hoteleve të regjistruara. Hapat në vazhdim për përdoruesin janë zgjedhja e hotelit dhe e opsionit të vlerësimit ku mund të shkruajë mendimin për qëndrimin në atë hotel.

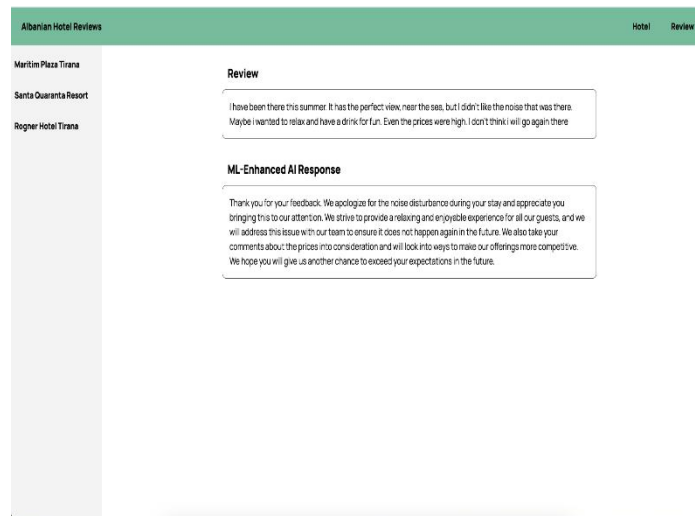


Figura 4.6 Përgjigje e gjeneruar nga Inteligjenca Artificiale

Automatikisht, pasi klienti shkruan vlerësimin, chatbot-i ynë përgjatit një përgjigje bazuar në rregullin e përcaktuar, informacionin e bazës së të dhënave dhe përgjigjen e simuluar nga GPT3.5 (Figura 4.6). Zgjat vetëm 3 sekonda, por është gjithashtu një njoftim për hotelin. Përmes modulit të integruar XAI, sistemi gjithashtu ofron një shpjegim bazë për përgjigjen e tij, duke e bërë të besueshme se po përgjigjet një njeri dhe jo një robot. Ky kombinim i arsytimit të shpjegueshëm dhe dialogut të ngjashëm me njeriun rrit besimin e mysafirëve dhe transparencën operacionale për operatorët e hoteleve. Përgjigjet e prodhuara jo vetëm që janë relevante në kontekst dhe të përshtatshme, por gjithashtu

zvogëlojnë nevojën që stafi njerëzor të trajtojë manualisht përgjigjet e përsëritura, duke kursyer kohë dhe duke siguruar qëndrueshmëri në zërin e markës, reagimet, duke kursyer kohë dhe duke siguruar qëndrueshmëri në zërin e markës.

4.3. Metodologjia

Materialet e nevojshme gjatë këtij hulumtimi përfshijnë disa biblioteka me funksionalitete të nevojshme. Duke filluar me Pandas ku tiparet e mysafirëve do të ruhen në një Data Frame, pastaj algoritmi Random Forest Classifier i tipit Supervised Learning, i cili përdoret kryesisht për klasifikim (madje ka edhe një version regresioni - RandomForestRegressor) por në rastin aktual, klasifikimi i përdoruesit është i mjaftueshëm. Për të shpjeguar se si vendos një model i Machine Learning, SHAP (SHapley Additive exPlanations) përdoret si një metodë e Explainable AI (XAI). Gjatë ndërtimit të disa funksionaliteteve, ishte e nevojshme të trajtoheshin vargjet numerike dhe dalja SHAP, e cila u arrit duke importuar bibliotekën NumPy (Numerical Python). NumPy përdoret gjerësisht për llogaritjet numerike dhe në rastet e punës me të dhëna në formën e matricave dhe vektorëve. Së fundmi, importohet biblioteka OpenAI, e cila shërben për të krijuar një lidhje të drejtpërdrejtë nga kodi ynë me OpenAI API në mënyrë që modelet e saj të mund të përdoren, siç specifikohet në GPT 3.5. Në këtë artikull, u bë një përpjekje për të gjeneruar tekst nga modeli GPT 3.5 duke dërguar një pyetje pak më formale sesa teksti i shkruar si vlerësim nga klienti.

```
FEATURES = ["loyalty_status", "stay_count", "room_available", "severity_high"]
def evaluate_user(user_dict):
    df = pd.DataFrame([user_dict])[FEATURES]
    prob = float(model.predict_proba(df)[0, 1][0])
    try:
        exp = explainer(df)
        shap_vals = np.array(exp.values)[0]
    except TypeError:
        vals = explainer.shap_values(df)
        if isinstance(vals, list):
            arr = vals[-1]
        else:
            arr = vals
        shap_vals = arr[0]
    explanation = dict(zip(FEATURES, shap_vals))
    return prob, explanation
```

Figura 4.7 Funkzioni për vlerësimin e përdoruesit

Funksioni `evaluate_user`, Figura IV.7, merr si të dhëna hyrëse të dhënat e mëparshme të përdoruesit dhe kthen vlerën e probabilitetit të llogaritur nga modeli parashikues dhe shpjegon parashikimin duke përdorur vlerat SHAP dhe nxjerr në pah karakteristikat që kanë ndikuar në parashikim. Versioni i vjetër i SHAP merret gjithashtu në konsideratë ku kthehet një listë ose një varg. Elementi i fundit i listës merret `vals[-1]` që shpesh është klasa pozitive përndryshe `arrivals`. Për të ruajtur informacionin për përdoruesin e parë specifikohet `arr[0]`. Funkzioni kthen probabilitetin dhe shpjegimin.

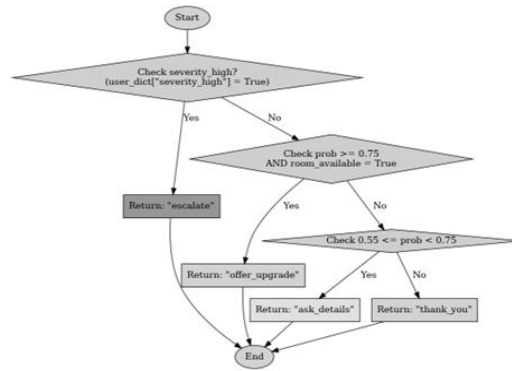


Figura 4.8 Paraqitja skematike e funksionit vendimarrës

Funksioni `decide_action` merr si parametra hyrës; probabilitetin e llogaritur në funksionin `evaluate_user` dhe informacionin e përdoruesit. Fluksi skematik i punës dhe vendimet e këtij funksioni përshkruhen nga Figura IV.8. Në këtë funksion, rëndësia ka parametri `severity_high` kur është i Vërtetë, gjë që injoron probabilitetin dhe e përshkallëzon menjëherë rastin në mbështetje më të lartë, që do të thotë se emergjencat kanë përparësi më të lartë se parashikimet. Nëse ka dhoma të lira dhe probabiliteti është $\geq 75\%$, sistemi i ofron klientit një përmirësim dhome, dhe ky rast është një kombinim i vërtetë i besimit të modelit me kufizimin e botës reale. Nëse probabiliteti është $\geq 55\%$ dhe $< 75\%$, klientit i kërkohen më shumë detaje për të ndërmarrë veprime për të shmangur vendimet e gabuara sepse vlerësohet se modeli nuk është i sigurt. Në rastet e një probabiliteti $< 55\%$, sistemi falënderon klientin dhe e mbyll rastin. Kërkesat efektive mund të përmirësojnë saktësinë, qëndrueshmërinë dhe rëndësinë në lidhje me informacionin e kërkuar nga GPT 3.5, kështu që për këtë arsye formulimi i kërkesave ka shumë rëndësi. Ka një punë shkencore në rritje për të udhëhequr dizajnin sistematik të kërkesave, veçanërisht për përdorime të specializuara si arsyetimi klinik ose edukimi. Ndërtimi i kërkesës u krye bazuar në disa nga strategjitë e specifikuar në artikull (White et al., 2023), ku fillimisht u përcaktua konteksti dhe fusha që përfshin pyetja. Zbërthimi i kërkesës u krye duke e ndarë atë në dy nënpika kryesore, lloji i veprimit dhe vlerësimi i klientit. Normalisht, këto dy informacione u vendosën si variabla që mund të ndryshoheshin nga klienti. Kreativiteti u përfshi gjithashtu si një element strategjik, ku metoda e përgjigjes u kërkuar përmes inxhinierisë së shpejtë si një kombinim i profesionalizmit dhe mirësjelljes.

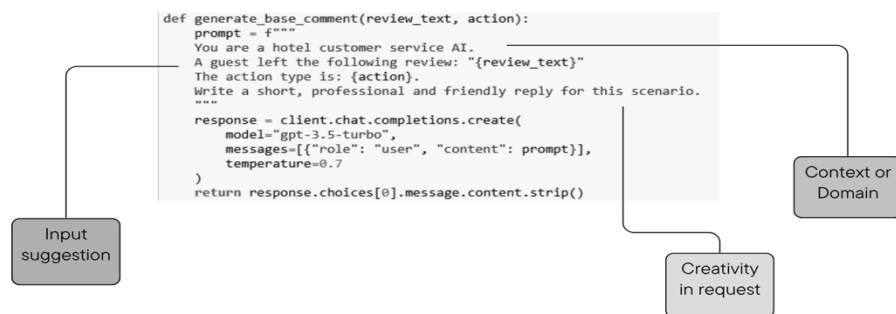


Figura 4.9 Grupimet e funksionit për gjenerimet bazë

Skematikisht, përshkrimi i funksionit `generate_base_comment` është dhënë në figurën IV.9, ku dallohen tre strategjitë e përdorura në formulimin e kërkesës. Bazuar në llojin e veprimit dhe shpjegimin, u vendos të modifikohet teksti i gjeneruar nga IA duke shtuar mesazhe të personalizuar. Për të siguruar që vlerat trajtohen në mënyrë të qëndrueshme kur llogariten, ishte i nevojshëm një funksion i dobishëm si `_scalar(x)` për konvertimin e `x` në një varg NumPy. Ky funksion u përdor brenda `modify_comment()`. Elementet hyrëse specifikohen `base_comment`, i cili është teksti i gjeneruar nga IA, veprimi si një varg që përfaqëson llojin e veprimit dhe shpjegimi i cili specifikon vlerat ose faktorët e rëndësishëm. Brenda funksionit, kërkohen dy faktorët më të rëndësishëm, ku artikujt e fjalorit të shpjegimit renditen sipas rëndësisë absolute, që do të thotë më i larti i pari dhe përdorimi i `_scalar` për të siguruar që çdo vlerë e rëndësishme trajtohet në mënyrë të qëndrueshme. Ne jemi të interesuar të ruajmë emrin e faktorëve që janë çelësi i fjalorit. Së fundmi, përmes kushteve në varësi të vlerës së veprimit, arrihet personalizimi, ku arsyeja shoqërohet me tekstin formal sipas kushteve. Duke peshuar faktorët më të rëndësishëm nga shpjegimi i modelit të inteligjencës artificiale, kjo veçori gjeneron një përgjigje të përmirësuar dhe të personalizuar. Veprimet kategorizohen si më poshtë:

- Oferta për optimizim shërbimi
- Kërkesa për detaje shtesë
- Përcjellje për trajtim të mëtejshëm
- Shprehje vlerësimi

```
def process_review(user_dict, review_text):
    prob, explanation = evaluate_user(user_dict)
    action = decide_action(prob, user_dict)
    base_comment = generate_base_comment(review_text, action)
    final_comment = modify_comment(base_comment, action, explanation)

    return {
        "probability": prob,
        "action": action,
        "final_comment": final_comment,
        "explanation": explanation
    }
```

Figura 4.10 Funksioni rishikimit të procesit

Funksioni përfundimtar është `process_review` i cili merr informacionin e përdoruesit dhe `review_text` si input dhe thërret të gjitha funksionet e mësipërme brenda tij. Funksioni i kthimit është një strukturë të dhënash që përfshin probabilitetin, veprimin, `final_comment` dhe shpjegimin. Kodi për këtë funksion tregohet në figurën IV.10.

4.4. Analiza e vlerësimit të modelit dhe të dhënave

Në shikim të parë, rezultati dukej shumë konsistent, bazuar në përgjigjen e marrë. Megjithatë, vlerësimi i saktësisë dhe rezultati F1 për të dhënat nxori në dritë gjetje të reja. Të dhënat përbëhen nga 500 të dhëna në formatin `xlsx`, dhe të dhënat u ndanë, 80% për trajnim dhe 20% për testim. Nuk ka vlera që mungojnë dhe formati është i lexueshëm. Objektivi (në këtë rast `severity_high`) është në formë binare (0/1) dhe është i përshtatshëm për klasifikim. Seti i trenit ka 400 rreshta ku klasa 0 ka 314 raste dhe klasa 1 ka 86 raste. Ndërkohë, seti i testimit ka 100 rreshta ku klasa 0 ka 78 raste dhe klasa 1 ka 22 raste. Në

këtë mënyrë, është arritur balanca e të dhënave të të dhënave dhe saktësia në këtë rast është 0.82. Ndërkohë, për të llogaritur F1-Score, mjafton të importohet ky funksion nga sklearn.metrics ku parametrat janë etiketat reale, parashikimet e modelit dhe mesatarja e ponderuar e cila ndryshon sipas madhësisë së klasës (domethënë, klasa me më shumë shembuj ka një peshë më të madhe). Kjo është më e vërtetë kur të dhënat janë të pabalancuara. Meqenëse F1-score është mesatarja harmonike e saktësisë dhe kujtesës, dhe në këtë rast vlerësohet si 74%, do të thotë që modeli ka një ekuilibër të mirë midis këtyre të dyjave, megjithëse ka ende vend për përmirësim. I njëjti vlerësim u krye për secilën klasë; dhe u arritën rezultate të ndryshme dhe të papritura. Classification_report u përdor për të nxjerrë në pah secilën klasë.

- Saktësia (kolona precision) - Tregon se sa nga rastet e parashikuara si klasë janë të asaj klase.
- Recall - Tregon se sa nga rastet reale të një klase u gjetën nga modeli.
- Rezultati F1-score - Është mesatarja e harmonizuar e saktësisë dhe recall.
- Suporti - Është numri aktual i mostrave për secilën klasë në të dhëna.

Sipas raportit të klasifikimit në Figurën IV.11, vihet re se modeli funksionon shumë mirë për klasën dominuese 0 (me 83 mostra) ku 83% e rasteve rezultuan 0.

```
from sklearn.metrics import classification_report
print(classification_report(y_test, y_pred))
```

	precision	recall	f1-score	support
0	0.83	0.99	0.90	83
1	0.00	0.00	0.00	17
accuracy			0.82	100
macro avg	0.41	0.49	0.45	100
weighted avg	0.69	0.82	0.75	100

Figura 4.11 Raporti klasifikimit të dy klasave

Modeli kapi 99% të rasteve reale. Ndërkohë, situata është e ndryshme për klasën 1 (me 17 mostra) ku përfundimi nënkupton se modeli nuk ka arritur të identifikojë asnjë rast nga kjo klasë dhe kjo për shkak të të dhënave të pabalancuara, d.m.th. më shumë 0 mostra se 1. Modeli ka mësuar të favorizojë klasën dominuese 0 sepse është më e lehtë të arrihet saktësi e lartë. Për të përmirësuar pak situatën, u përdorën disa teknika balancimi si Smote (Teknika e Mbi-mostrave Sintetike të Minoriteteve), nën-mostra dhe balancimi i peshës së klasave.

SMOTE				
	precision	recall	f1-score	support
0	0.800	0.564	0.662	78
1	0.244	0.500	0.328	22
accuracy			0.550	100
macro avg	0.522	0.532	0.495	100
weighted avg	0.678	0.550	0.588	100
AUC: 0.49067599067599077				
UnderSampler				
	precision	recall	f1-score	support
0	0.717	0.423	0.532	78
1	0.167	0.409	0.237	22
accuracy			0.420	100
macro avg	0.442	0.416	0.385	100
weighted avg	0.596	0.420	0.467	100
AUC: 0.4463869463869464				
Class_weight='balanced'				
	precision	recall	f1-score	support
0	0.757	0.679	0.716	78
1	0.167	0.227	0.192	22
accuracy			0.580	100
macro avg	0.462	0.453	0.454	100
weighted avg	0.627	0.580	0.601	100
AUC: 0.486013986013986				

Figura 4.12 Rezultatet e klasifikimit pas aplikimit të teknikave balancuese

Çdo teknikë u aplikua në secilën klasë, dhe rezultatet tregohen në Figurën IV.12. Teknika SMOTE ka një rënie në saktësi prej 55%. Ajo krijon mostra sintetike për klasën minore (klasa 1). Klasa 0 ka saktësi të lartë për kujtesë të ulët, kështu që modeli gjen disa raste, por lë shumë raste të pazbuluara. Ndërsa përsëri për klasën 1 saktësia është shumë e ulët 0.24 dhe kujtesa 0.50 që do të thotë gjysmën e kohës arrin të gjejë raste të klasës 1, por shpesh bën gabime. Kjo teknikë ka performancë më të mirë në kujtesë, por është saktësi shumë e ulët. Teknika e nën-mostrimit ka një rënie në saktësi prej 42% dhe gjatë kësaj procedure shumë mostra hiqen nga klasa 0 për të balancuar të dhënat. Për klasën 0 saktësia është 0.71 dhe kujtesa është 0.42 (shumë e ulët), ndërsa për klasën 1 saktësia është 0.17 dhe kujtesa 0.41 (arri të gjejë raste, por shumë e pasaktë). Saktësia ka rezultuar shumë e ulët sepse kemi humbur shumë informacion nga klasa 0. Si rezultat, kjo teknikë nuk po funksionon mirë sepse humbasin shumë të dhëna nga klasa e madhe. Kjo teknikë humbi informacion dhe uli performancën. Teknika e Peshës së Klasës (e balancuar) rezulton në një saktësi prej 48%. Modeli u jep më shumë peshë gabimeve në klasën 1. Sa i përket klasës 0, saktësia vlerësohet në 0.76 dhe rikujtesa është 0.68, që duket të jetë një situatë pak më e mirë se teknikat e mësipërme. Për klasën 1, saktësia është 0.17 dhe rikujtesa është 0.23, që është përsëri shumë e dobët, por pak më e mirë se nën-mostra. Meqenëse rezultati F1 i kësaj teknike është afërsisht 60%, mund të themi se modeli ka pak më shumë ekuilibër, megjithatë, klasa 1 ende nuk po mëson mirë. Kjo teknikë është një kompromis, por nuk i përmirësoi rezultatet për klasën 1.

4.5. Teknikat e balancimit

Në këtë studim propozohet një strukturë chatbot-i e shpjgueshme e drejtuar nga IA në fushën e mikpritjes. Ai përfshinte një kombinim të modelit të ansamblit Random Forest, duke përdorur shumë pemë vendimesh, metodën e shpjgimit matematik SHAP, duke u përqendruar dhe duke nxjerrë në pah kryesisht disa atribuide të veçorive, dhe një model të madh gjuhësor si GPT-3.5 për të ofruar përgjigje bazuar në vlerësimet e klientëve. Sistemi bazohet kryesisht në transparencë, në ndërgjegjësimin për kontekstin për vendimmarrje, në

balancimin e automatizimit bazuar në rregulla, në mënyrë që rastet kritike të mund të prioritizohen. Duke integruar interpretueshmërinë në model dhe në përgjigjen e klientit, qasja rrit besimin, auditueshmërinë dhe personalizimin. Hulumtimet e ardhshme do të përqendrohen në vlerësime në shkallë më të gjerë, krahasimin me teknikat alternative XAI dhe validimin e përqendruar të përdoruesit për të forcuar përvetësimin në mjediset e mikpritjes në botën reale. Një çështje e rëndësishme konsiston në nxjerrjen në pah të rëndësisë së adresimit të çekuilibrit të të dhënave në modelet parashikuese për sektorin e mikpritjes.

Tabela 4.1 Vlerësimi empirik e teknikave balancuese për modelin klasifikues të Chatbot-it

<i>Teknikat balancuese</i>	<i>Saktësia (%)</i>	<i>F1-Score (Klasa minoritare)</i>	<i>F1-Score (Klasa maxhoritare)</i>	<i>AUC (Area Under Curve)</i>	<i>Interpretimi</i>
<i>SMOTE (Synthetic Minority Oversampling Technique)</i>	55	0.328	—	0.49	EkUILIBër i përmirësuar midis klasave, por ndashmëri e kufizuar; mbi-mostrimi sintetik rrit saktësinë, por prapëseprapë shkakton klasifikim të gabuar.
<i>Random Undersampling</i>	42	0.237	—	—	Metoda më pak efektive: rënie e performancës për shkak të humbjes së të dhënave nga zvogëlimi i klasës së shumicës.
<i>Class Weight = 'balanced'</i>	58	0.192	0.716	—	Rezultatet e përgjithshme më të qëndrueshme, por performanca e dobët e klasës së pakicave, tregojnë potencial për përshtatje të ndjeshme ndaj kostos.

Rezultatet e analizës paraqiten në Tabelën 4.1 dhe sqarojnë se teknikat ekzistuese për të zgjidhur problemin e disbalancës në të mësuarit automatik, të tilla si SMOTE, random undersampling dhe class weight balance, ndikojnë ndjeshëm në performancën e klasës së pakicës, por asnjëra nuk dha rezultate të qëndrueshme. Përdorimi i SMOTE arriti një ekuilibër më të mirë midis klasave krahasuar me metodat e tjera. Saktësia ishte 55% dhe Rezultati F1 i klasës së pakicës ishte 0.328, madje edhe ato me një AUC (Zona nën kurbë) relativisht të ulët (0.49). Kjo tregon se rritja artificiale e të dhënave mund të përmirësojë saktësinë, por nuk arrin të ofrojë diferencim të qartë midis klasave. Metoda e Nën-mostrave rezultoi më pak efektive, me një saktësi prej 42% dhe një rezultat të ulët F1 (0.237) për klasën pozitive, duke sugjeruar humbje informacioni gjatë reduktimit të mostrës. Ndërkohë, përdorimi i peshës së klasës= 'i balancuar' dha rezultate disi më të qëndrueshme, duke arritur një saktësi prej 58% dhe një rezultat F1 më të lartë për klasën kryesore (0.716), por me dobësi të dukshme në klasën e pakicës (rezultati F1 0.192). Kjo sugjeron që për përmirësim të ndjeshëm kërkohen qasje më të avancuara, siç janë mësimi në ansambël, algoritmet e ndjeshme ndaj kostos ose kombinimet hibride të metodave

ekzistuese. Studimi konfirmon se çekuilibri i të dhënave është një faktor kritik për arritjen e rezultateve të besueshme në aplikimet e inteligjencës artificiale në mikpritje dhe turizëm. Në përgjithësi, këto rezultate tregojnë se çekuilibri i të dhënave mbetet një sfidë kritike në zhvillimin e sistemeve të IA-së që janë të besueshme për sektorin e mikpritjes. Për të arritur një performancë më të fuqishme parashikuese, kërkime të mëtejshme duhet të përqendrohen në adresimin e këtij çekuilibri përmes teknikave të përparuara të rimodelimit (SMOTE me TOMEK Links ose ADASYN) dhe duke mbledhur të dhëna shtesë (p.sh., duke rritur të dhënat) për klasën e pakicës. Për më tepër, eksperimentimi me algoritme më të fuqishme, siç janë XGBoost, LightGBM ose rrjetet nervore me humbje fokale, mund të përmirësojë performancën parashikuese. Së fundmi, një vlerësim më i gjerë duke përdorur metrika validimi specifike për domenin, duke përfshirë kurbat e kujtesës së saktësisë dhe indeksin e kënaqësisë së klientit, do të jetë thelbësor për të konfirmuar zbatueshmërinë dhe shkallëzueshmërinë në botën reale të kornizës së propozuar.

4.6. Analiza konkluzioneve

Ky punim propozon një strukturë të shpjegueshme të chatbot-it të drejtuar nga IA në fushën e mikpritjes. Ai përfshinte një kombinim të modelit RandomForest, shpjegimin e SHAP bazuar në disa attribute të veçorive dhe përgjigjet e gjeneruara nga GPT-3.5. Fokusi i sistemit është transparenca, i vetëdijshëm për kontekstin për vendimmarrje, balancimi i automatizimit bazuar në rregulla në mënyrë që rastet kritike të mund të prioritizohen. Duke integruar interpretueshmërinë në model dhe në përgjigjen e klientit, qasja rrit besimin, auditueshmërinë dhe personalizimin. Hulumtimet e ardhshme do të përqendrohen në vlerësime në shkallë më të gjerë, krahasimin me teknikat alternative XAI dhe validimin e përqendruar të përdoruesit për të forcuar përvetësimin në mjediset e mikpritjes në botën reale. Një çështje e rëndësishme është se ky studim theksoi rëndësinë e adresimit të çekuilibrit të të dhënave në modelet parashikuese për sektorin e mikpritjes. Analizat treguan se teknikat ekzistuese si SMOTE, nën-mostra dhe peshat e klasave ndikojnë ndjeshëm në performancën e klasës së pakicës, por asnjëra nuk ofroi rezultate të qëndrueshme. Kjo sugjeron që qasje më të avancuara si të mësuarit në ansambl, algoritmet e ndjeshme ndaj kostos ose kombinimet hibride të metodave ekzistuese janë të nevojshme për përmirësim të ndjeshëm. Studimi konfirmon se çekuilibri i të dhënave është një faktor kritik për arritjen e rezultateve të besueshme në aplikimet e inteligjencës artificiale në mikpritje dhe turizëm. Një nga problemet e hasura gjatë kësaj pune janë të dhënat e pabalancuara, ku klasa e pakicës është e nën-përfaqësuar krahasuar me klasën e shumicës. Kjo pabarazi rezulton në një performancë të ulët për klasën e pakicës, dhe kjo pasqyrohet më së miri në rezultatet e rikujtimit dhe F1. Megjithatë u aplikuan teknikat e balancimit, një rregullim i pjesshëm i problemit ndodhi përsëri, dhe vlerat e AUC (Sipërfaqja nën kurbë) mbeten afër parashikimit të rastësishëm. Një drejtim premtues për eksplorim në të ardhmen është adresimi i kësaj pabarazie përmes teknikave të përparuara të rimodelimit (SMOTE me TOMEK Links ose ADASYN) dhe duke mbledhur të dhëna shtesë (p.sh., duke rritur të dhënat) për klasën e pakicës. Për më tepër, eksperimentimi me algoritme më të fuqishme si XGBoost, LightGBM ose rrjete nervore me humbje fokale mund të rrisë performancën parashikuese. Së fundmi, një vlerësim më i gjerë me metrika validimi specifike për domenin do të jetë i nevojshëm për të siguruar që sistemi është i besueshëm dhe i zbatueshëm në skenarë të botës reale.

KAPITULLI V

Analiza e Sentimentit nga e Folura në Gjuhën Shqipe: Një Qasje e të Mësuarit të Thellë për Skenarë me Burime të Kufizuara dhe vendimarrja automatike

Zëri është ai që mbart informacion të fuqishëm paralinguistike, dhe në konsideratë merren intonacioni, lartësia e zërit, energjia që përcillet dhe timbri. Me anë të tyre zbulohen emocione përtej fjalëve. Këto sinjale akustike ndihmojnë në saktësimin e qëndrimit emocional nëse shoqërohen edhe me transkriptimin e tekstit. Njohja e emocioneve nga e folura mundëson ndërveprime empatike dhe natyrale me makinat. Kjo vlen në rastet kur kërkohet reagimi ndaj nuacave emocionale për të përmirësuar ndjeshëm përvojën e përdoruesit në aplikime si asistentë virtualë, platforma edukative, kujdesi shëndetësor dhe shërbimi ndaj klientit [134]. Në rastet kur e folura mbetet e vetmja mënyrë për zbulimin e gjendjes emocionale kjo sjell fuqizimin e sistemeve të reagimit në kohë reale, e mjetëve të aksesueshmërisë dhe platformave të automatizuara që të përgjigjen në mënyrë adekuate pa qënë e nevojshme të transkriptohet teksti [134]. Saktësia më e lartë dhe kuptim më i nuancuar arrihet me kombinimin e të dhënave të nxjerra nga emocionet e të folurës dhe analizës sentimentale të tekstit. Sistemet multimodale që kombinojnë veçoritë akustike me modelet NLP shfaqin dukshëm performancë më të mirë në njohjen e gjendjes emocionale [135]. Teknikat e të mësuarit të vetë-mbikëqyrur (supervised learning), si HuBERT, wav2vec ose UniSpeech-SAT kanë shfaqur cilësi të mira dhe të forta në kapjen e shenjave të emocionit dhe sentimentit drejtpërdrejtë nga e folura edhe në rastet e mungesës së tekstit. Njohja e emocioneve nga e folura (SER-Speech Emotion Recognition) është shfaqur si një teknologji thelbësore me shtrirje të gjërë në shumë fusha. Në shëndetësi, SER mundëson zbulimin e hershëm të problemeve të shëndetit mendor dhe monitorimin e mirëqënies psikologjike, veçanërisht duke marrë parasysh rritjen e problemeve mendore gjatë pandemisë COVID-19 [136]. Teknologjia ka përmirësuar ndërveprimin mes njeriut dhe kompjuterit duke krijuar sisteme të Inteligjencës Artificiale më empatike dhe duke zhvilluar më tej shërbimet e Asistentëve Inteligjentë Virtualë Personal [137]. Në arsim, njohja e emocioneve nga e folura luan një rol të rëndësishëm për të kuptuar më shumë nga sjellja dhe angazhimi i nxënësve. Mësuesit mund të shfrytëzojnë analizën e sentimentit për të vlerësuar gjendjen emocionale të nxënësve gjatë aktiviteteve mësimore, duke i ndihmuar të përshtasin metodat për të maksimizuar nxënien dhe përmirësuar rezultatet arsimore [138]. Meqë proceset njohëse dhe të nxënies ndikohen nga emocionet, sistemet SER i mundësojnë mësuesve të mbështesin më mirë mirëqënien dhe suksesin akademik të nxënësve [139]. Këto aplikime dëshmojnë rëndësinë e SER në krijimin e mjediseve dixhitale më gjithëpërfshirëse dhe emocionalisht inteligjente, në sektorin e shëndetësisë, arsimit dhe teknologjisë.

5.1. Gjuha shqipe si një gjuhë me burime të kufizuara

Për gjuhën shqipe gjenden shumë pak tekste të etiketuara për të realizuar detyra si Analiza e Sentimentit, Njohja e Entiteteve të Emërtuara (NER-Named Entity Recognition), klasifikimi etj. Konstatohet relativisht një numër i vogël punimesh për sentimentin e gjuhës shqipe krahasuar me gjuhët e mëdha [140];[141]. Në identifikimin e gjuhës shqipe si gjuhë

me burime të kufizuara, përdorimi i Rrjeteve Neurale Artificiale (ANN-Artificial Neural Networks) dhe Rrjeteve Neurale Konvolucionale (CNN-Convolutional Neural Networks) ka treguar aftësi të mira në të nxënin e modeleve gjuhësore të shqipes [142]. Situata paraqitet më qartazi përmes rasteve të dokumentuara në artikuj shkencorë dhe përmbledhet si mëposhtë:

- Sipas artikullit [143] shprehet që për gjuhën shqipe, studiuesit shpesh duhet të ndërtojnë manualisht tokenizuesit, lematizuesit dhe pipeline sepse ekzistojnë shumë pak mjete të Përpunimit të Gjuhës Natyrore (NLP) si spaCY, Stanza, HuggingFace models të cilat kryesisht janë të ndërtuar për gjuhë të tjera [143].
- Në krahasim me gjuhëm angleze, burimet e etiketimit në shqip janë shumë të pakta psh: dataset-i ALBNER me rreth 900 fjali të etiketuara [144], Databaza e sentimentit ALBMoRe me 800 komente filmash [145] dhe dataset-i për klasifikimin e temave, AlbNews me 600 tituj të etiketuar [146]
- Gjuha shqipe ka dy dialekte kryesore, Gegë dhe Toskë, dhe mungesa e një përfaqësimi të balancuar midis tyre e vështirëson ndërtimin e një modeli të qëndrueshëm.

5.2. Kategorizimi i karakteristikave akustike për njohjen e emocioneve

Njohja e Emocionit të të Folurit (SER-Sentiment Emotion Recognition) mbështetet në nxjerrjen e karakteristikave informative që kapin si vetitë fizike të sinjaleve të të folurit ashtu edhe përfaqësimet semantike të nivelit të lartë të mësuar nga modelet e thella. Në përgjithësi, këto karakteristika mund të kategorizohen në:

1. Karakteristikat prozodike si lartësia e tonit (F_0), intensiteti/fortësia, shpejtësia e të folurit, pauza, kontura e tonit, janë elementë të rëndësishëm për identifikimin e emocioneve, kryesisht në raste zemërimi, gëzimi edhe frike.
2. Karakteristikat Spektrale që përshkruajnë përmbajtjen frekuecore të sinjalit. Këtu përmendim MFCC (Mel-Frequency Cepstral Coefficients) si më e përdorura nga SER (Sentiment Emotion Recognition). Rast tjetër është edhe LPC (Linear Predictive Coding) që kap strukturën e traktit vokal dhe përdoret në Emotion Recognition për rastet si zemërimi/ entuziazmi që kanë formantë më të lartë dhe energji më të madhe, apo trishtimi që ka energji më të ulët dhe artikulum më të butë. LPC janë shumë të ndjeshme ndaj ndryshimeve emocionale, Tonnnetz dhe përfaqësojnë timbrin, strukturën e frekuencës dhe modelet akustike të ngarkuara me emocione dhe në një studim këto janë renditur midis karakteristikave kryesore të nxjerra për njohjen e emocioneve të të folurit (Rezapour Mashhadi dhe Osei-Bonsu, 2023). Ndërsa në një studim tjetër, kombinimi i MFCC-së, Kromës, spektrogramit Mel, Tonnnetz dhe kontrastit spektral rrit saktësinë e klasifikimit të emocioneve me 93% në RAVDESS. (Gondohanindijo et al., 2023) Karakteristikat spektrale janë shumë të dobishme në përdorim kryesisht për Machine Learning dhe Deep Learning.
3. Karakteristikat perceptuese (psikoakustike) që lidhen me atë që perceptohet nga njerëzit. Bazohet në standardet ITU PEAQ (vlerësimi perceptues i njësisë ndërkombëtare të telekomunikacionit për cilësinë e audios) ku karakteristika të tilla si zhurma e pjesshme, raporti i ndryshimit emocional me maskën, ndryshimi i zarfit, përmbajtja harmonike dhe shfaqja kohore e blloqeve emocionale shërbejnë për të identifikuar këtë perceptim nga njerëzit (Sezgin et al., 2012).

4. Qasjet e Bashkimit me Shumë Karakteristika që përfshijnë kombinimin e disa llojeve të karakteristikave për të arritur performancën më të mirë të mundshme. Studimet specifikojnë se përdorimi i karakteristikave të llojeve të ndryshme përmirëson klasifikimin e emocioneve me rreth 95% duke përdorur ANN të Rregulluar Bayesian (Gondohanindijo et al., 2023). Një studim tjetër që përdor të dhënat RAVDESS (Ryerson Audio-Visual Database of Emotional Speech and Song) vepron në audio të parapërpunuar, ku inkuadrimi, heqja e heshtjes etj. mund të jenë bërë për të krijuar një të dhënë njohjeje emocionesh (Colunga-Rodriguez et al., 2025)
5. Integrimet e Thella Audio, të cilat janë qasje moderne që shfrytëzojnë atë që mësohet nga rrjetet nervore. Modele si L3-Net dhe VGGish, mund të kapin semantikë emocionale të nivelit të lartë në muzikë pa karakteristika të krijuara me dorë (Koh dhe Dubnov, 2021).
6. Studimet për të përdorur integrimet e modelit ASR (Njohja Automatike e të Folurit) për njohjen e emocioneve dhe nganjëherë tejkalojnë grupet klasike të karakteristikave si eGeMAPS (Grupi i Parametrave Akustikë Minimalistë të Gjenezës i zgjeruar, i cili është një grup standard i karakteristikave akustike për kryerjen e analizës së të folurit dhe emocioneve) (Tits et al., 2018).

5.3. Metodologjia e nxjerrjes së karakteristikave për njohjen e emocioneve nga audio

Për të analizuar një audio, ekzistojnë disa elementë të rëndësishëm si amplituda ose energjia që identifikon se sa i lartë është tingulli, shkalla e kalimit zero (ZCR) që do të thotë se sa shpesh sinjali kalon zeron për të identifikuar tingujt e zëshëm kundrejt atyre pa zë, dhe Energjia Kohëshkurtër që mat energjinë në fragmente të shkurtra dhe ndihmon në zbulimin e aktivitetit të të folurit. Këto karakteristika nxirren direkt nga forma e valës. Normalisht, fokusi është në karakteristikat që lidhen me zbulimin e emocioneve. Në fakt, karakteristikat janë përfaqësime numerike të një sinjali audio. Ato shërbejnë si një input për modelin e Mësimi Automatik ose Mësimi të Thellë. Karakteristikat më të njohura akustike janë:

- Toni: Një ton i mesëm ose i lartë pasqyron emocione të tilla si zemërimi, eksitimi, etj., ndërsa një ton më i ulët sugjeron trishtim, mërzitje, etj. [147].
- Energjia: identifikon intensitetin e emocioneve të tilla si trishtimi, gëzimi që kanë nivele të larta energjie.
- Karakteristika spektrale: Mel-spektrogramet, MFCC-të, qendra spektrale, fluksi ofrojnë sinjale të përmbajtjes timbrale dhe frekuencore [148]
- Karakteristika prozodike: Kohëzgjatja, ritmi i të folurit, pauzat, modelet e intonacionit ndihmojnë në dallimin e gjendjeve emocionale. Rëndësia qëndron në kontekstin "si u tha dhe jo çfarë u tha".
- Metrikat e cilësisë së zërit: përmenden dridhja e zërit, frymëmarrja apo edhe elementë të tjerë që i shtojnë kuptim emocional zërit me nuanca të ndryshme.

Për nxjerrjen e karakteristikave ekzistojnë disa metoda tradicionale dhe moderne. Metodat e zgjedhura në këtë artikull janë MFCC (Koeficientët Cepstral të frekuencës Mel) si një metodë tradicionale dhe HuBERT (Hudden-Bnit ERT), wav2vec 2.0 wavLM si metoda moderne. HuBERT është një Model i të Folurit i Vetë-Mbikëqyrur (SSM) që nxjerr ngulitje

nga audio. Në këtë rast, segmenti audio i jepet modelit dhe nxirret një vektor me dimensione fikse (768-1024). Me këtë metodë, karakteristikat komplekse nxirren pa pasur nevojë për etiketim. Ndërkohë, metoda tradicionale MFCC nxjerr 13-40 koeficientë për kornizë, mesatarja dhe devijimi standard llogariten për secilin segment.

5.4. Përgatitja datasetit

Zëri i regjistruar është i një personi femër. Audio bazohet në leximin e disa fjalive neutrale ose edhe me shënime emocionale pozitive dhe negative. Janë krijuar tre klasa bazë, ku këto audio janë grupuar sipas emocioneve: pozitive, negative dhe neutrale. Për secilën klasë janë krijuar rreth 300 mostra audio me intervale kohore nga 3 deri në 9 sekonda. Për të llogaritur kohëzgjatjen e audios përdoret raporti midis numrit të frame-ve dhe sample rate-it, siç paraqitet në formulën mëposhtë:

$$\text{Duration (sec)} = \frac{\text{Number of frames}}{\text{Sample rate}} \quad (5.1)$$

Llogaritja e statistikave përshkruese përfshin minimumin si një algoritëm linear që iteron nëpër të gjitha vlerat dhe ruan më të voglën, maksimumin që gjen vlerën më të madhe, medianën që rendit listën (zakonisht me kompleksitet $O(n \log n)$) dhe zgjedh vlerën e mesme, si edhe kohëzgjatjen totale për secilën klasë, e cila llogaritet duke pjesëtuar shumën e kohës së të gjitha audiove me 3600 për të marrë totalin në orë.

Tabela 5.1 Statistikë vlerësuese e dataset-it audio

	<i>Pozitive</i>	<i>Negative</i>	<i>Neutrale</i>
<i>Kohëzgjatja Max</i>	10.14 s	9.25 s	11.10 s
<i>Kohëzgjatja Min</i>	2.82 s	3.09 s	2.55 s
<i>Kohëzgjatja Mes</i>	4.99 s	5.59 s	6.59 s
<i>Kohëzgjatja total</i>	0.44 h	0.46 h	0.54 h
<i>Nr i skedarëve</i>	295	296	295

Tabela 5.1 paraqet statistikën vlerësuese të dataset-it audio për tre klasat e sentimentit: pozitive, negative dhe neutrale. Rezultatet tregojnë se kohëzgjatja mesatare e skedarëve audio varion nga 4.99 s në klasën pozitive deri në 6.59 s në klasën neutrale. Gjithashtu, numri i skedarëve për secilën klasë është pothuajse i balancuar, me rreth 295–296 skedarë për kategori, çka ndihmon në reduktimin e paragjykimeve gjatë trajnimit të modeleve të inteligjencës artificiale. Dataset-i përmban gjithashtu variacion në kohëzgjatjen minimale dhe maksimale të audios, duke reflektuar diversitetin e të dhënave zanore të përdorura në analizën e sentimentit.

Mjeti i përdorur për regjistrimin e audios është një laptop me pajisje Mikrofonit AMD Audio. Tekstet që u regjistruan janë përmbajtje që gjenden rastësisht në internet. Cilësimet hyrëse të mikrofonit janë në formatin 2 Kanale, 16 bit, 48000 Hz (Cilësia e DVD-së). Rëndësia nuk qëndron vetëm në regjistrimin e audios, por edhe në krijimin e një seti të dhënash excel me formatin csv (vlera të ndara me presje). Fushat e këtij formati janë path, si vendndodhja e audiove, dhe etiketat, si një etiketë për secilën audio. Përgatitja e setit të të dhënave përfshin gjithashtu procesin e heqjes së zhurmës, normalizimit ose edhe

krijimin e segmenteve dhe kornizave. Para nxjerrjes së veçorive, është e rëndësishme që audiot të jenë 16 kHz mono si një format i përshtatshëm për t'u përdorur në setet e të dhënave të Njohjes Automatike të të Folurit dhe Njohjes së Emocioneve të të Folurit.

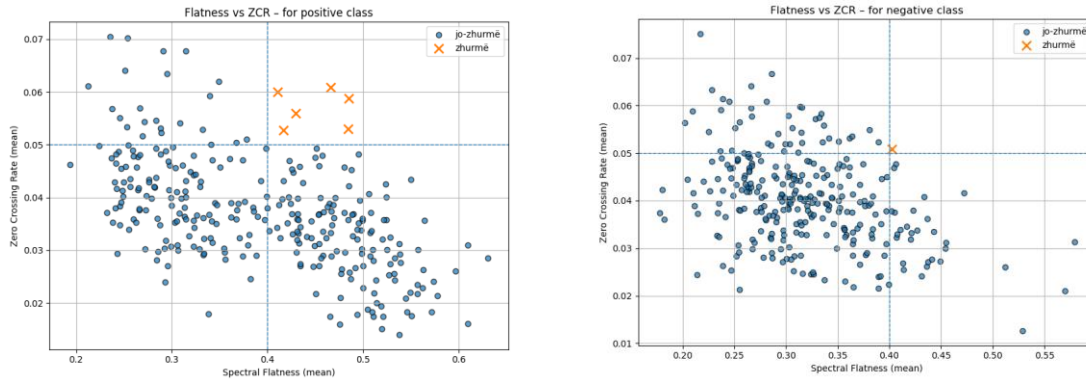


Figura 5.1 Vlera mesatare për S.Flatness dhe ZCR Pozitive dhe Negative

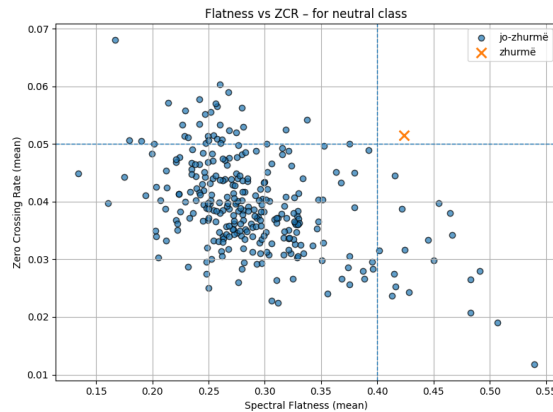


Figura 5.2 Vlera mesatare për S.Flatness dhe ZCR Neutrale

Për klasifikimin e ndjesive nga audioja, duhet të merren në konsideratë shumë karakteristika akustike. Figura 5.2 paraqet mesataren e dy karakteristikave, siç janë Sheshësia Spektrale dhe Shkalla e Kalimit Zero, bazuar në të dhënat e të dhënave të krijuara. E para dallon tonalitetin nga zhurma, ndërsa e dyta dallon zërin e qetë nga ai i zënë.

Në klasën negative, sipas sheshësisë spektrale, shumica e pikave janë midis 0.20-0.50 me strukturë të moderuar, ndërsa ZCR ka vlera midis 0.02-0.07, të cilat janë tipike për të folurit normal. Siç mund të shihet, vetëm një pikë ka kaluar pragun dhe konsiderohet si zhurmë. Si përfundim, klasa negative është relativisht e pastër dhe ka shumë pak segmente që duken si zhurmë. Në klasën Pozitive, sheshësia spektrale varion nga 0.25 në mbi 0.5, që do të thotë se ka më shumë zhurmë. Ndërkohë, ZCR ka një shpërndarje më të gjerë ku disa pika tejkalojnë 0.06. Kjo klasë ka më shumë zhurmë të perceptuar krahasuar me të tjerat ose më saktë ka më shumë elementë të turbullt akustikisht, kështu që përfshin segmente jo të pastra. Në klasën neutrale, shumica e pikave të rrafshësisë spektrale kanë vlera rreth 0.25-0.35 dhe kjo është e ngjashme me klasën negative, ndërsa ZCR ka më pak variacion. Kjo klasë ka audio më të pastër krahasuar me klasën negative. Si rezultat, klasa

që ka më shumë zhurmë akustike me më shumë segmente ku rrafshësia është >0.45 dhe ZCR më e lartë se 0.06 është klasa pozitive. Llogaritja e ZCR (shkalla e kalimit zero) kryhet duke përdorur një algoritëm linear $O(n)$ dhe vlera e tij është më e lartë kur ka zhurmë në audio, ndërsa llogaritja e rrafshësisë spektrale bazohet në algoritmin $O(n \log n)$. Sa më afër vlerës 1, aq më "i sheshtë" (si zhurmë) është spektri; sa më afër vlerës 0, aq më "tonal/i strukturuar" (fjalim, muzikë).

Figura 5.3 tregon Shkallën mesatare të Kalimit Zero (ZCR) për tre klasa audio: Pozitive, Negative dhe Neutrale. Boshti X (0–1): koha e normalizuar (nga fillimi deri në fund të audios, pavarësisht nga gjatësia absolute). Boshti Y (0–0.08): ZCR mesatare, e cila tregon se sa shpesh sinjali audio kalon zeron në secilën njësi të kohës. Tre vijat (blu, portokalli, jeshile): ZCR mesatare për klasat Pozitive, Negative dhe Neutrale. Për të interpretuar grafikun, ne e ndajmë boshtin Y (që është ZCR) në tre segmente.

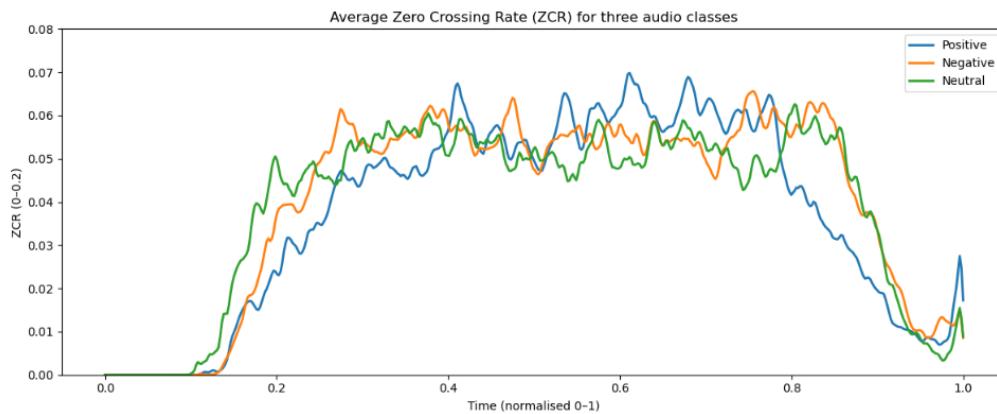


Figura 5.3 Mesatarja e ZCR për klasat e audios

Segmenti i parë (0–0.2) Ku vërehet një rritje e shpejtë e ZCR për të gjitha klasat. Ky fenomen tregon se në fillim të audios ka shumë tranzicione të shpejta të sinjalit (fillimi i fjalës, frymëmarrja, tinguj të mprehtë). Klasa Neutrale (e gjelbër) tenton të ketë një ZCR pak më të lartë, gjë që mund të sugjerojë tinguj më të shpërndarë ose më pak të qëndrueshëm në hyrje.

Segmenti qendror (0.2–0.8) Linjat janë më të qëndrueshme dhe i afrohen vlerave mesatarisht 0.05–0.07. Ky stabilizim sugjeron pjesë vokale më të qëndrueshme, me dallime të vogla midis klasave. Dallimet midis klasave janë më të dukshme rreth 0.3–0.5: Pozitivja dhe Negativja arrijnë pika më të larta se Neutralja, të cilat mund të lidhen me intensitet më të madh ose tinguj më energjikë në këto emocione.

Segmenti i fundit (0.8–1.0) Është vërejtur se ZCR bie ndjeshëm për të gjitha klasat, duke treguar një rënie në aktivitetin e sinjalit (mbarimi ose pauza e fjalës). Pozitivja bie më herët dhe më fort, gjë që mund të sugjerojë mbyllje më të butë të fjalës krahasuar me Neutralja dhe Negativja. Të tre klasat ndjekin një profil të ngjashëm (ngritje fillestare → pllajë qendrore → rënie në fund), i cili është karakteristikë e strukturës së përgjithshme të të folurit. Neutralja ka ZCR pak më të lartë në fillim, gjë që mund të pasqyrojë tinguj të shpërndarë ose fillime më të shkurtra fjalësh. ZCR-të pozitive dhe negative tregojnë luhatje më të theksuara gjatë segmentit qendror, i cili shoqërohet me prodhim vokal më energjik dhe ndoshta ndryshime në frekuencë për shkak të ngarkesës emocionale. ZCR-ja pozitive tregon një rënie më të theksuar në fund, duke sugjeruar qartësi më të madhe në fund të të folurit.

5.4.1. Strategjia e etiketimit të emocioneve të audios

Në analizën e ndjenjës dhe emocioneve nga audioja, cilësia dhe mënyra e etiketimit të të dhënave është po aq e rëndësishme sa zgjedhja e modelit ose karakteristikave akustike. Një debat i rëndësishëm është: a duhet të bëhet etiketimi në nivelin origjinal audio (niveli i shprehjes) apo pas përpunimit paraprak/segmentimit (niveli i segmentit ose i kuadrit)? Shumica e veprave klasike bëjnë etiketim në nivelin e shprehjes (p.sh., një regjistrim i plotë etiketohet “pozitiv”, “negativ” ose “neutral”), dhe më pas përdoret përpunimi paraprak për të segmentuar sinjalin në njësi më të vogla për trajnimin e modelit. Segmentet shpesh trashëgojnë etiketën e shprehjes, megjithëse ato mund të mos përmbajnë të gjithë ngarkesën emocionale. Ky proces është i thjeshtë dhe efikas, por krijon një rrezik të zhurmës së etiketimit kur segmentet janë heterogjene. Prandaj, për segmentet audio, një algoritëm metrik i distancës, siç është distanca Euklidiane, u përdor për të vlerësuar heterogjenitetin e segmentimeve.

$$d(x, y) = \sqrt{\sum_i (x_i - y_i)^2} \quad (5.2)$$

Formula e distancës Euklidiane

Ky algoritëm u aplikua në secilën klasë duke i grupuar ato sipas emrit të skedarit. Bazuar në vlerat ku në rastet $\text{mean_distance} \geq 280$ identifikohen rastet ku segmentet audio janë heterogjene, ndërsa kur kjo distancë është ≥ 100 por < 280 kategorizohen segmentet që kanë pjesë të ngjashme brenda grupit dhe së fundmi nëse $\text{mean_distance} < 100$, segmentet audio vlerësohen si me ngjashmëri të lartë. Tabela 5.2 paraqet vlerat e mean_distance sipas klasave.

Tabela 5.2 Distanca mesatare ndërmjet segmenteve

	<i>Pozitive</i>	<i>Negative</i>	<i>Neutrale</i>
<i>Heterogjen</i>	0	0	0
<i>Disa ngjashmëri</i>	85	100	95
<i>Të ngjashëm</i>	207	196	199
Totali	293	296	294

Një artikull propozon përdorimin e një metode për etiketat pseudo-emocionale në nivel kornize (grupim) dhe më pas përdorimin e etiketave për të rafinuar modelin. Kjo metodë adreson drejtpërdrejt faktin se "jo të gjitha frekuencat/kornizat" brenda një shprehjeje përfaqësojnë emocionalitetin e plotë [149]. Pra, procesi përfshin etiketimin e shprehjes përpara se audioja të përpunohet paraprakisht.

5.4.2. Paraprosesimi

Përpunimi audio kalon nëpër disa faza. E para është normalizimi i amplitudës (0-1) ose [-

1,1]. Me normalizim të amplitudës nënkuptohet që sinjali të bëhet i një shkalle uniforme dhe të jetë i sigurt për përpunim, p.sh. i unifikuar në vëllim. Duke aplikuar transformimin linear $f(x)=2x+1$ rezulton se nëse $x \in [-1,1]$ atëherë $f(x) \in [0,1]$. Pastaj është krijimi i segmenteve 1-2 sekonda. Segmentet 1-2 sekonda përdoren sepse ato ofrojnë kontekst të mjaftueshëm emocional (intonacion, ritëm, energji) pa përzier shumë gjendje të ndryshme në të njëjtin shembull dhe pa mbingarkuar kujtesën gjatë trajnimit. Meqenëse emocioni ndryshon ngadalë në të folur, fragmente prej disa sekondash janë të nevojshme për të kapur modelet prozodike dhe kjo theksohet në veprat që kërkojnë varësi të gjata kohore në SER [150]. Segmentet më të gjata rrisin "zhurmën e etiketës" ku emocione të ndryshme mund të ekzistojnë brenda të njëjtës shprehje, dhe fragmentimi i tyre në pjesë më të shkurtra zvogëlon problemin e kësaj ngatërrese [151]. Heqja e zhurmës dhe heshtjes së panevojshme (dezhurmimi ose heqja e heshtjes) sepse mjediset reale janë të zhurmshme dhe kjo çon në një rënie të ndjeshme të [152]. Zhurma shtrembëron karakteristikat (ndryshon energjinë në breza, maja, etj.) kështu që klasifikuesit ngatërrohen [153]. Së fundmi, ekziston shkalla e marrjes së mostrave, e cila në rastin e HuBERT është 16 kHz, ndërsa për MFCC është 16-44 kHz. Shkalla e marrjes së mostrave do të thotë se sa mostra marrim për sekondë nga vala e zërit. Është e rëndësishme sepse specifikon se sa detaje kapim në valën e zërit brenda një sekonde. Si rezultat, me normalizim krijohet një stabilitet numerik dhe amplituda të krahasueshme, me heqjen e heshtjes dhe dezhurmimin mund të arrihet një SNR më i lartë efektiv dhe më pak zhurmë etiketimi në segmente. Me anë të resampling ofrohet një mbështetje e qëndrueshme kohore-frekuencë për veçorinë MFCC/SSL dhe me anë të segmentimit për grumbullimin do të fiksohen forma që mundësojnë etiketimin e segmentit ose të etiketës së kornizës.

5.5. Ekstradimi i karakteristikave me MFCC dhe HuBERT

Duke krahasuar metodat MFCC dhe HuBERT për nxjerrjen e veçorive audio, u zbulua se HuBERT zakonisht tejkalon MFCC me saktësi 1.33-10.46% më të lartë në të gjitha aplikacionet, megjithëse HuBERT kërkon më shumë burime llogaritëse dhe ka raste kur avantazhet e tij mund të ulen në varësi të domenit. Në një studim mbi zbulimin e dizartrisë dhe klasifikimin e ashpërsisë [154] raportoi se veçoritë HuBERT dhanë saktësi zbulimi 1.33–2.86% më të lartë dhe saktësi klasifikimi të ashpërsisë 6.54–10.46% më të madhe sesa MFCC dhe linjat e tjera bazë. Gjatë Njohjes Automatike të të Folurit në artikullin nga Kumar u vu re se performanca e HuBERT ndryshon në varësi të domenit, theksit dhe gjuhës, ndërsa bankat e filtrave Mel ndonjëherë performojnë më mirë [155]. Studime të tjera që vlerësojnë HuBERT në njohjen e të folurit [156] dhe identifikimin e emocioneve [157] raportuan metrika të favorshme, të tilla si një shkallë gabimi në fjalë deri në 4.6-7.6% dhe një shkallë njohjeje prej 82.6%, megjithëse nuk u dha asnjë krahasim i drejtpërdrejtë i MFCC. Në skenarin brenda gjuhësor, modelet SSL (p.sh. HuBERT) arritën deri në 34.4% përmirësim në testimin e karakteristikave artikuluese (AF) mbi MFCC, ndërsa në skenarin ndër-gjuhësor përmirësimi është 26.7% mbi MFCC [158]

5.6. Klasifikues neural i lehtë për analizën e emocionit në audio

Një klasifikues nervor i lehtë është një rrjet nervor i thjeshtë që përdoret për të klasifikuar veçoritë e nxjerra me MFCC, HuBERT, etj. dhe ofron një ekuilibër midis performancës

dhe kostos. Modele të tilla preferohen kur grupet e të dhënave janë të vogla ose kur kërkohet përfundim në kohë reale.

5.6.1. Rrjeti Feed Forward me dy shtresa

MLP (Multi-Layer Perceptron) është një arkitekturë e rrjeteve nervore artificiale e përbërë nga disa shtresa neuronesh, ku neuronet e çdo shtrese lidhen plotësisht me neuronet e shtresës pasuese, duke formuar një strukturë të dendur lidhjesh. Fillimisht, ekziston një shtresë që merr karakteristika (p.sh. vektorë nga MFCC ose ngulitje nga HuBERT) e cila konsiderohet shtresa hyrëse, pastaj ekzistojnë një ose më shumë shtresa të fshehura që ruajnë informacion jolinear të të dhënave dhe shtresa dalëse ku ekziston një shtresë që prodhon klasa (pozitive, negative, neutrale). Rrjeti me Feed Forward me Dy Shtresa është një rast i thjeshtë i MLP që përmban 2 shtresa të fshehura. MLP me Feed Forward me Dy Shtresa përdoret për klasifikimin e karakteristikave audio [159]. Klasifikuesi MLP importohet nga scikit-learn, dhe madhësia e shtresës së fshehur është vendosur në 128 që është numri i neuroneve. Në analizën e ndjenjës audio, funksioni i aktivizimit ReLU (Njësia Lineare e Rregulluar) përdoret për të siguruar stabilitet në shtresat e fshehura. Forma e funksionit të aktivizimit ReLU është si më poshtë figura 5.4:

:

$$f(x) = \max(0, x)$$

Figura 5.4 Funksioni aktivizimit ReLU

Funksioni i aktivizimit ReLU kalon vlera pozitive dhe zero në vlera negative. Është shumë i shpejtë dhe ndihmon në shmangien e problemit të zhdukjes së gradientëve. Pra, në shtresën e fshehur kemi Shtresë të Dendur e cila matematikisht përfaqësohet nga formula e mëposhtme në figurën 5.5 së bashku me funksionin e aktivizimit ReLU.

$$h = f(Wx + b)$$

Figura 5.5 Funksioni shtresës Dense

5.6.2. Modeli klasifikues Dy shtresor Feed Forward për MFCC dhe HuBERT

Fillimisht, çdo skedar csv u ngarkua veçmas, për MFCC dhe HuBERT në mënyrë që të krahasohej performanca në varësi të asaj që merret nga këta nxjerrës të veçorive. Ndërkohë, fokusi është në stratifikimin për ndarjen e të dhënave (p.sh., ndarja e trajnimit/validimit/testit ose validimi i kryqëzuar). Stratifikimi është procesi që siguron që përqindja e klasave të mbahet e njëjtë në secilën ndarje si në të dhënat origjinale. Nëse stratifikimi nuk zbatohet, atëherë ndarja do të jetë e rastësishme dhe mund të prodhohen shpërndarje jo-proporcionale. Ai garanton që secila klasë të përfaqësohet në të gjitha ndarjet, duke siguruar besueshmëri në saktësinë, F1, metrikat e Matricës së Konfuzionit dhe duke shmangur rrezikun e "ndarjeve të paragjykuara" ku një model mund të duket më

i mirë ose më i keq thjesht sepse një klasë nuk është testuar. Stratifikimi i përdorur bazohet në Algoritmin e Mostrës së Stratifikuar të Rastësishme. Hapi tjetër është aplikimi i klasifikuesit me dy shtresa Feed-Forward dhe vlerësimi i performancës që përfshin saktësinë, Makro-F1, Raportin e Klasifikimit dhe Matricën e Konfuzionit.

Tabela 5.3 Raporti Klasifikues

<i>Modeli klasifikues</i>	<i>Precizioni</i>			<i>Madhësia F1</i>		
	Negativ	Neutral	Pozitive	Negativ	Neutral	Pozitive
<i>MFCC-MLP(128)</i>	0.681	0.717	0.681	0.671	0.74	0.657
<i>HuBERT-MLP(128)</i>	0.652	0.754	0.677	0.668	0.754	0.655

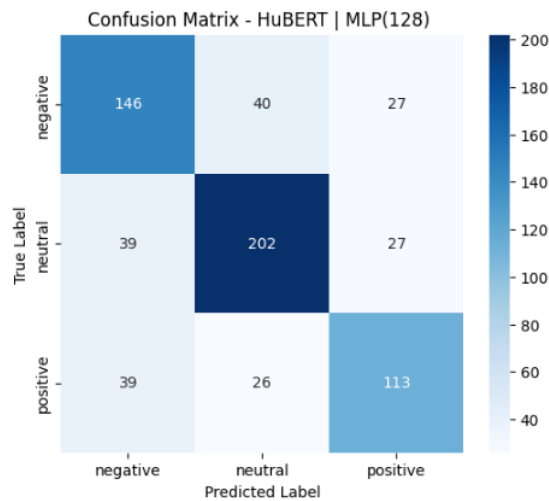


Figura 5.6 Matrica e konfuzionit për MFCC për klasifikuesin MLP

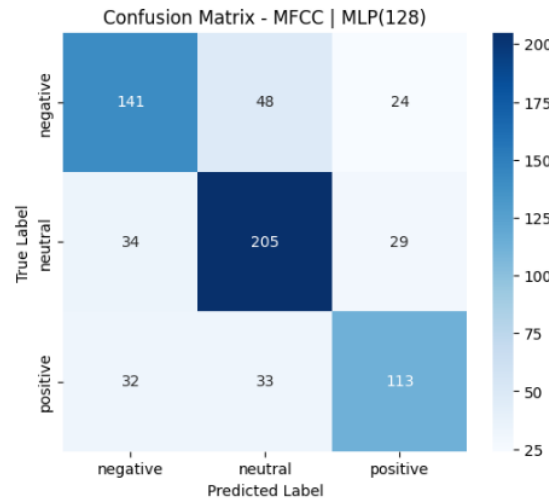


Figura 5.7 Matrica e konfuzionit për HuBERT për klasifikuesin MLP

Matrica e konfuzionit përfaqëson një teknikë të përdorur gjerësisht për analizimin dhe vlerësimin e saktësisë së modeleve të klasifikimit. duke ofruar informacion më të detajuar sesa një metrikë e vetme siç është saktësia ose rezultati F1. Ajo tregon se ku ngatërrohen klasat, gjë që nuk është e dukshme nga saktësia. Paraqitet si një matricë ku rreshtat përfaqësojnë etiketat reale (e vërteta bazë), ndërsa kolonat përfaqësojnë etiketat e parashikuara nga modeli. Elementet diagonale (p.sh. "negative → negative") tregojnë numrin e rasteve të klasifikuara saktë, ndërsa elementët jashtë diagonaleve tregojnë gabime specifike të modelit (p.sh. "pozitive" të klasifikuara si "neutrale"). Ky mjet është i dobishëm për të kuptuar shpërndarjen e gabimeve dhe për të identifikuar klasat që modeli shpesh i ngatërrohen. Në rastin e rezultateve tona, të dy konfigurimet (MFCC + MLP, HuBERT + MLP) tregojnë se klasa neutrale ka saktësinë më të lartë, ndërsa klasat negative dhe pozitive shpesh ngatërrohen me njëra-tjetrën. Kjo pritet në analizën e ndjenjës nga audio, ku vija ndarëse midis emocioneve negative dhe pozitive është shpesh e hollë dhe varet nga nuancat e intonacionit [160]. Për më tepër, krahasimi i matricës së konfuzionit midis modeleve tregon qartë se përdorimi i MLP në ngulitje HuBERT jep një shpërndarje më të balancuar të klasifikimeve të sakta, duke zvogëluar gabimet ndërklasore. Kjo nxjerr në pah avantazhin e përdorimit të modeleve të vetë-mbikëqyrura krahasuar me tiparet tradicionale të MFCC [161]

Tabela 5.4 Saktësia dhe Madhësia F1

		Saktësia		Macro F1	
		Test	Valid	Test	Valid
<i>HuBERT</i>	MLP(128)	0.7	0.728	0.692	0.721
<i>MFCC</i>	MLP(128)	0.697	0.729	0.689	0.727

Rezultatet eksperimentale, të paraqitura në Tabelën 5.4, treguan se modeli HuBERT + MLP(128) arriti një vlerë macro-F1 rreth 0.69 në dataset-in e testimit, ndërsa modelet e bazuara vetëm në MFCC ose klasifikues linearë (LogReg) shfaqën performancë disi më të ulët (0.64–0.68). Ky nivel performance është në përputhje me literaturën ekzistuese mbi njohjen e sentimentit/emocioneve nga e folura në gjuhë me burime të kufizuara, ku modelet tipike të bazuara vetëm në MFCC raportojnë vlera F1 në intervalin 0.60–0.70.[162] [160]. Ndërkohë, përdorimi i embedding-ve të gjeneruara nga modele të vetë-mbikëqyrura (self-supervised) si HuBERT ose wav2vec 2.0 është treguar se përmirëson ndjeshëm performancën krahasuar me MFCC tradicionale, por pa arritur nivelet e performancës së gjuhëve që disponojnë dataset-e shumë të mëdha [161]. Diferenca e vogël midis performancës në validim dhe testim (rreth 2–3%) sugjeron se overfitting nuk përbën problem dhe se modeli është në gjendje të përgjithësojë në mënyrë të qëndrueshme. Këto rezultate tregojnë se, megjithëse F1=0.69 nuk është shumë i lartë krahasuar me skenarët “high-resource”, ai përfaqëson një performancë të besueshme dhe të pritshme për analizën e sentimentit nga audio në kontekste ku dataset-i është i kufizuar dhe gjuha ka pak burime, si në rastin e gjuhës shqipe.

5.6.3. LSTM dydrejtimore (Bi-LSTM) me Shtresën e vëmendjes

Bi-directional LSTM (Bi-LSTM) është një nga arkitekturat më të përdorura sot në njohjen

e emocioneve nga e folura (SER) dhe në analizën e sentimentit. [163]. Bi-LSTM është një variant i RNN-ve (Recurrent Neural Networks) që kap varësitë afatgjata në sekuenca kohore. Këto sekuenca përpunohen në të dy drejtimet (fillim $\rightarrow$ fund dhe fund $\rightarrow$ fillim). Prandaj, ai ka rëndësi të veçantë kur informacioni merret si nga e kaluara ashtu edhe nga e ardhmja brenda segmentit për të interpretuar intonacionin dhe ritmin e audios.[164]. Problemi me LSTM-të është se jo të gjitha frame-t e MFCC-ve (ose embedding-et HuBERT) kanë të njëjtën peshë për detyrën e klasifikimit. Për këtë arsye merret në konsideratë një Attention Layer (shtresë vëmendjeje), e cila i mundëson modelit të “përqendrojë” vëmendjen te frame-t që përmbajnë më shumë informacion, si për shembull ato ku ndodhin ndryshime të forta të energjisë ose tonalitetit (pitch), të cilat tregojnë emocion. Matematikisht, mekanizmi i vëmendjes llogarit një shpërndarje peshash mbi të gjitha hapat kohorë dhe kombinon gjendjet e fshehura në një përfaqësim të peshuar [165]. Kombinimi BiLSTM + Attention është bërë standard në shumë aplikime të njohjes së sentimentit/emocioneve nga e folura. BiLSTM kap kontekstin në të dy drejtimet, ndërsa mekanizmi Attention përzgjedh frame-t më informuese, duke rritur kështu interpretueshmërinë dhe duke dalluar se cilat pjesë të audios ndikojnë më shumë në vendimin e modelit [166]. BiLSTM me Attention është një arkitekturë hibride që përmirëson ndjeshëm performancën krahasuar me BiLSTM të thjeshtë, veçanërisht në detyrat ku dinamikat kohore dhe karakteristikat akustike luajnë një rol kyç. Figura 5.8 paraqet një imazh perceptual të matricës me frame kohore dhe vektor.

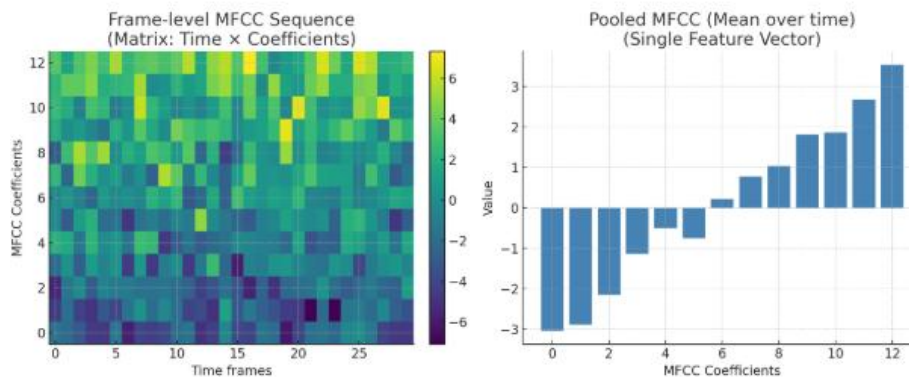


Figura 5.8 Perceptimi vizual i matricës me frame kohore dhe vektor

Ne kemi shumë segmente audio dhe secili prej tyre konvertohet në një sekuençë frame-sh. Çdo frame përmban 5 koeficientë (MFCC_0..MFCC_4) $\rightarrow$ një vektor me dimension $D=5$, por segmentet audio kanë gjatësi të ndryshme ($T_i \neq T_j$). Modelet neuronale zakonisht presin tensorë me gjatësi fikse. Renditja kohore realizohet sipas path-it dhe frame-it për të ruajtur rendin temporal të frame-ve (shmanget “përzjerja” kohore). Për çdo skedar p, merret një matricë $(T,5)$ — T frame dhe 5 karakteristika për secilin frame. Label-i caktohet vetëm një herë për të gjithë segmentin audio (supozim: i gjithë segmenti ka të njëjtin sentiment). Label encoding realizohet duke i konvertuar etiketat nga formë tekstuale në indekse (p.sh. neutral=0, positive=1, negative=2). Kjo është e nevojshme për funksionin e humbjes në klasifikim (p.sh. CrossEntropy). Gjithashtu përcaktohen datatype-et e sakta për

PyTorch (float32 për input-et dhe long për label-et).

Prandaj, gjatë gjithë këtij punimi, e folura është trajtuar si një seri kohore shumëdimensionale, ku çdo frame përfaqëson një vektor karakteristikash akustike. Kur është e nevojshme të realizohet klasifikimi i sentimentit (negative, neutral, positive), kërkohet një vektor fiks për marrjen e vendimit, por sekuencat kanë gjatësi të ndryshme dhe për këtë arsye përdoret modeli BiLSTM + pooling. Ky model synon të lexojë të gjithë sekuencën e karakteristikave (MFCC) duke marrë kontekst si nga e kaluara ashtu edhe nga e ardhmja. Sekuenca me gjatësi T përmbledhet në një vektor të vetëm fiks (pooling në kohë) dhe ky vektor projektohet në hapësirën e klasave përmes një shtrese linear + softmax për të gjeneruar probabilitetet.

Arkitektura përfshin Input-in $X \in \mathbb{R}^{B \times T \times D}$, ku B është madhësia e batch-it, T numri i frame-ve (i cili ndryshon midis segmenteve audio) dhe $D = feat_dim$, ku në rastin tonë janë përdorur 5 karakteristika për çdo frame. Pas input-it ndodhet LSTM, i cili përpunon sekuencën në dy drejtime: majtas $\rightarrow$ djathtas dhe djathtas $\rightarrow$ majtas. Output-i për çdo hap kohor është:

$$H_t \in \mathbb{R}^{2H} \quad (5.3)$$

(që përfaqëson bashkimin e dy drejtimeve), ndërsa output-i total është:

$$H_{all} \in \mathbb{R}^{B \times T \times 2H} \quad (5.4)$$

Pooling në kohë (mesatarizimi) jep:

$$A \in \mathbb{R}^{B \times 2H} \quad (5.5)$$

duke realizuar mesatarizimin përgjatë boshtit kohor dhe duke e bërë përfaqësimin invariant

ndaj gjatësisë së sekuencës.

$$\bar{h} = \frac{1}{T} \sum_{t=1}^T H_t \quad (5.6)$$

Pjesa e fundit në këtë arkitekturë është klasifikuesi linear + softmax:

$$Z = WH + b \quad (5.7)$$

Softmax gjeneron:

$$p(y | x) = \text{softmax}(Z) \quad (5.6)$$

Ndërsa funksioni i humbjes trajtohet me Cross-Entropy:

$$L = \text{CrossEntropy}(Z, y) \quad (5.7)$$



Figura 5.9 Kurba e humbjes gjatë trajnimit dhe validimit – BiLSTM + Attention + MFCC

Për të kuptuar sjelljen e modelit gjatë trajnimit, përdoret paraqitja vizuale përmes grafikut të Figura 5.9. Grafiku paraqet dy kurba: Blu (train loss) dhe Portokalli (validation loss). Vërehet se train loss ka një ulje të vazhdueshme nga $\approx 1.08 \rightarrow \approx 0.92$ dhe kjo tregon se modeli po mëson mirë mbi të dhënat e trajnimit. Ndërkohë, validation loss fillon rreth 1.06, bie lehtësisht në ≈ 1.04 dhe më pas mbetet pothuajse e sheshtë, çka do të thotë se modeli nuk po përmirësohet shumë mbi të dhënat e validimit.

Nga grafiku konkludohet se nuk ka shenja të forta të overfitting-ut, sepse në rastin klasik të overfitting-ut kurba e train loss do të vazhdonte të binte ndërsa validation loss do të fillonte të rritej qartë. Në rastin tonë, train loss bie, ndërsa validation loss mbetet e qëndrueshme, duke sugjeruar se modeli mund të ketë arritur kapacitetin maksimal me këto parametra.

Duke qenë se diferenca midis train dhe validation loss është e vogël (~ 0.1), themi se përmirësimi ka stagnuar, pra modeli vazhdon të mësojë gjatë trajnimit por nuk përfiton më në validim. Kjo sugjeron që arkitektura (BiLSTM 128, average pooling) ose karakteristikat e përdorura nuk ofrojnë më informacion të mjaftueshëm për të reduktuar më

tejt validation loss.

Meqenëse validation loss nuk përmirësohet pas disa epoch-esh, ndalimi i trajnimit është i justifikuar, pasi kursen kohë dhe shmang over-training. Për optimizim, është e nevojshme të rritet numri i karakteristikave, p.sh. $13-20 \text{ MFCC} + \Delta + \Delta\Delta$. Gjithashtu, në vend të average pooling mund të provohet attention-pooling ose përdorimi i gjendjes së fundit të fshehur (last hidden state). Kapaciteti i modelit mund të rritet duke përdorur më shumë njësi, p.sh. 256, ose dy shtresa BiLSTM.

Në përfundim, bazuar në rezultatet, mund të themi se modeli po mëson, sepse train loss po zvogëlohet vazhdimisht dhe modeli po përshtatet me të dhënat e trajnimit. Megjithatë, performanca në validim mbetet problematike, pasi modeli nuk po përmirëson aftësinë e tij të përgjithësimit mbi të dhëna të reja. Vlerësohet se nuk ka shenja të dukshme të mbipërshtatjes, sepse validation loss nuk rritet ndjeshëm, por kemi një nënpërshtatje të pjesshëm, pasi modeli nuk arrin të mësojë një strukturë për të reduktuar ndjeshëm validation loss. Trajnimi i modelit ndalet pas epokës së 3–4 kur përmirësimi në validim ndalon. Për të kursyer burime kompjuterike dhe për të parandaluar mbitrajnimin, përdoren teknika të “early stopping”.

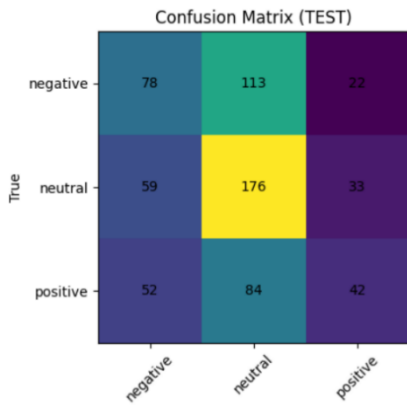


Figura 5.10 Matrica konfuzionit për Bi-LSTM me Shtresën e Vëmendjes përpara class weight loss

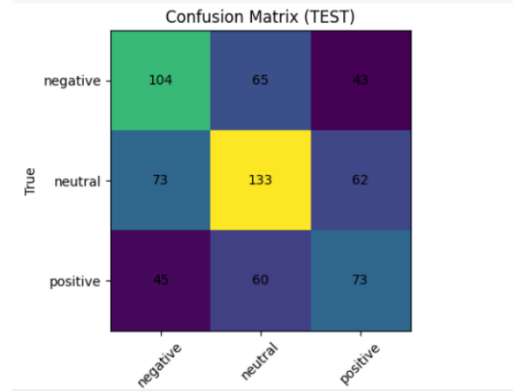


Figura 5.11 Matrica e konfuzionit për Bi-LSTM me Shtresën e Vëmendjes pas weight class loss

Në dataset-in e testimit (held-out test set), modeli BiLSTM-Attention arrin 45% accuracy (macro-F1 0.41). Matrica e konfuzionit (Figura 5.10 dhe 5.11) zbulon një paragjykim sistematik drejt klasës neutrale (373 parashikime), me shumë raste të klasifikimit të gabuar si neutrale si për klasën negative ashtu edhe për atë pozitive. Performanca sipas klasave është e pabarabartë (F1: neutral 0.55, negative 0.39, positive 0.31), duke treguar aftësi të kufizuar dalluese të përfaqësimit aktual MFCC dhe/ose problem të balancimit të klasave. Përmirësimet duhet të fokusohen në cilësinë e karakteristikave dhe në reduktimin e problemit të pabalancimit të klasave. Megjithatë, duke vendosur class weight, vërehet një përmirësim në matricën e konfuzionit, sepse favorizohet klasa minoritare. Me cross-entropy standard, humbja për një mostër të klasës c është:

$$\ell = -\log p_{\theta}(y=c | x). \quad (5.8)$$

Ndërkohë, me peshat e klasave w_c , formula është:

$$\ell = -w_c \log p_{\theta}(y=c | x). \quad (5.9)$$

Sipas raportit të klasifikimit, rezultatet tregojnë më shumë gjasa për underfitting dhe një pabalancim/paqartësi të lehtë të etiketave, jo domosdoshmërisht overfitting, pasi performanca e përgjithshme është mesatare.

Tabela 5.5 Raporti klasifikues

Modelet klasifikues	Precizioni			Madhësia F1		
	Negative	Neutrale	Pozitive	Negative	Neutrale	Pozitive
MFCC-BiLSTM+Vëmendje)	0.41	0.47	0.43	0.39	0.55	0.31
Weight class loss kundrejt entropisë standarte	0.41	0.52	0.46	0.42	0.50	0.37
HuBERT-BiLSTM+Vëmendje)	0.55	0.70	0.59	0.57	0.64	0.61

Tabela 5.5 krahason performancat e tre qasjeve të ndryshme për klasifikimin e modeleve të njohjes së emocioneve bazuar në zë, duke përdorur metrika standarde si Precision dhe F1-score për tre klasat (positive, negative dhe neutral). Modeli MFCC-BiLSTM+Attention paraqet rezultate modeste për precision (0.41–0.47) dhe një F1-score të ulët, kryesisht për klasën Positive (0.31). Kjo tregon se, megjithëse MFCC ofron një bazë solide karakteristikash akustike, ekzistojnë kufizime në dallimin e emocioneve më komplekse. Për të përmirësuar rezultatet me MFCC, u konsiderua një rast tjetër. Në vend të entropisë standarde, u përdor weight class loss dhe rezultatet treguan përmirësime të lehta krahasuar me entropinë standarde, kryesisht për klasat neutral dhe positive (precision 0.46, F1-score 0.37). Weight class loss ndihmoi në balancimin e dataset-it, megjithëse efekti mbeti i kufizuar. Rasti i tretë është HuBERT-BiLSTM+Attention dhe ky kombinim paraqet performancën më të mirë, me precision 0.70 për klasën Neutral dhe F1-score deri në 0.64 për Neutral dhe 0.61 për Positive. Kjo reflekton më së miri se modelet e vetë-mbikëqyrura (self-supervised) si HuBERT kapin informacion semantik dhe prosodik më të pasur krahasuar me modelet tradicionale. Përmirësimi u vu re veçanërisht për klasën Positive (Figura 5.12), e cila shpesh ishte më sfiduese për t'u identifikuar.

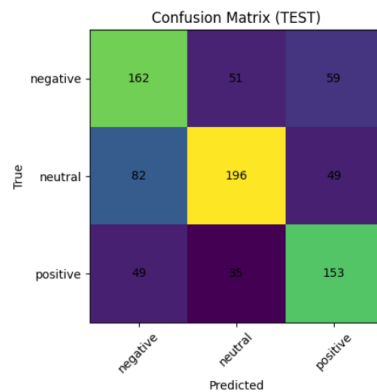


Figura 5.12 Matrica e konfuzionit për HuBERT+BiLSTM me shtresën e vëmendjes

KAPITULLI VI

Framework sugjerues për Vendimmarrjen Automatike bazuar në Binjakun Dixhital

6.1. Prezantimi i problemit

Ky kapitull paraqet kornizën e propozuar për realizimin e vendimmarrjes automatike në menaxhimin e cilësisë së ajrit të brendshëm, me fokus të veçantë në kontrollin e lagështirës. Qasja e propozuar integron të dhëna në kohë reale nga sensorët, teknika të inteligjencës artificiale për parashikim, dhe një model DT (Digital Twin) me qëllim kalimin nga një sistem reaktiv në një sistem proaktiv dhe inteligjent.

Në këtë kontekst, DT (Digital Twin) shërben si një përfaqësim virtual i ambientit fizik, duke mundësuar simulimin dhe optimizimin e vendimeve përpara zbatimit të tyre në sistemin real. Menaxhimi i lagështirës në ambientet e brendshme është një proces dinamik, i cili ndikohet nga faktorë të ndryshëm si temperatura, ventilimi dhe aktivitetet njerëzore. Sistemet tradicionale të kontrollit bazohen kryesisht në rregulla statike (p.sh., aktivizimi i pajisjes kur lagështira kalon një prag të caktuar), të cilat kanë kufizime të konsiderueshme. Në ndryshim nga qasjet tradicionale, ku vendimet merren në mënyrë reaktive bazuar në pragje fikse, ky punim synon të implementojë një sistem që:

- parashikon gjendjen e ardhshme të ambientit,
- simulon sjelljen e sistemit në një model virtual,
- dhe aktivizon automatikisht pajisjen për kontrollin e lagështirës duke njoftuar edhe individin.

Kjo qasje përbën një hap drejt sistemeve inteligjente të menaxhimit të ambientit të brendshëm, të bazuara në të dhëna dhe vendimmarrje automatike.

Lagështira është një element i rëndësishëm i cilësisë së ajrit sepse ndikon në zhvillimin e mikroorganizmave (myk, baktere), ndikon në cilësinë e frymarrjes dhe degradimin e materialeve. Në këtë kontekst sistemet tradicionale të kontrollit janë reaktive (veprojnë pasi problemi ndodh), jo-adaptive dhe pa sugjerim të inteligjencës vendimarrëse. Qasja e re e propozuar gjatë këtij rasti studimor është lidhja e DT me Inteligjencën Artificiale në ndihmesë të realizimit të Vendimarrjes Automatike.

6.2. Koncepti i Binjakut Dixhital

Koncepti i Digital Twin (DT) është shfaqur si një qasje moderne për lidhjen dhe ndërveprimin ndërmjet sistemeve fizike dhe përfaqësimeve të tyre dixhitale inteligjente. Termi u prezantua fillimisht nga Grieves, i cili e përshkroi DT-në si një kopje virtuale të një objekti fizik, e aftë të shoqërojë sistemin real gjatë gjithë ciklit të tij të jetës, duke mbështetur monitorimin, analizën dhe optimizimin në kohë reale. [167]. Më vonë, NASA e zgjeroi këtë koncept duke e përdorur për menaxhimin e sistemeve komplekse hapësinore, ku binjaku dixhital shërbente për simulimin e gjendjes së komponentëve dhe parashikimin e dështimeve të mundshme [168]

Në literaturën bashkëkohore, Binjaku Dixhital përkufizohet si një përfaqësim virtual

dinamik i një entiteti fizik ose procesi, i cili përditësohet vazhdimisht nga të dhëna reale përmes sensorëve, IoT (Internet of Things), databazave dhe rrjeteve të komunikimit, duke krijuar një lidhje të sinkronizuar ndërmjet botës fizike dhe asaj dixhitale [169]. Sipas Tao et al., një DT kombinon modelimin fizik, të dhënat operative dhe algoritmet inteligjente për të realizuar simulim, parashikim, diagnostikim dhe vendimmarrje në kohë reale [170]. Ndërsa Fuller et al. theksojnë se vlera kryesore e DT qëndrojnë në aftësinë për të mbështetur optimizimin operacional dhe autonominë e sistemeve përmes analizës së vazhdueshme të të dhënave [171].

Në këtë kuptim, DT konsiderohet një evolucion i sistemeve tradicionale të simulimit, pasi nuk kufizohet në analiza statike offline, por operon si një sistem kibernetik-fizik inteligjent me feedback të vazhdueshëm [169]. Për këtë arsye, përdorimi i tij është zgjeruar nga industria prodhuese drejt sektorëve si shëndetësia, energjia, qytetet inteligjente dhe menaxhimi i mjedisve të brendshme, ku kërkohet monitorim adaptiv dhe vendimmarrje automatike [171].

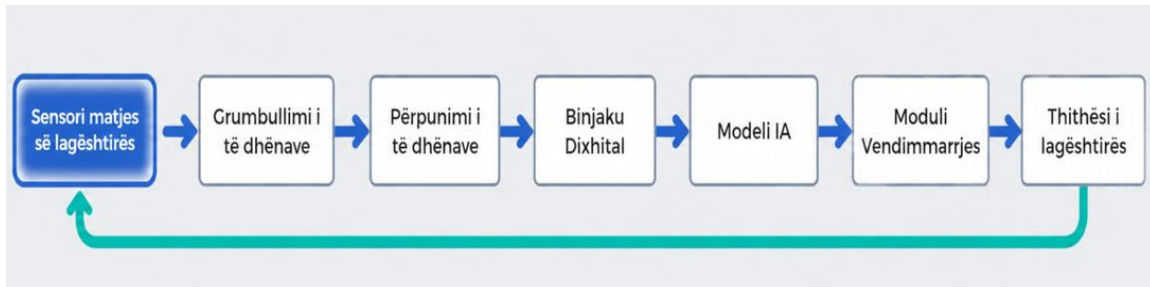


Figura 6.1 Framework-u i Binjakut Dixhital (DT) në ndihmesë të realizimit të vendimmarrjes

Në arkitekturën e propozuar, Figura 6.1, Binjaku Dixhital (DT) përfaqëson komponentin virtual inteligjent që lidh shtresën fizike me shtresën e Inteligjencës Artificiale dhe vendimmarrjes automatike. Roli i tij nuk kufizohet vetëm në paraqitjen virtuale të ambientit, por në sinkronizimin e vazhdueshëm të të dhënave reale të lagështirës me modelin virtual të sistemit. Pasi të dhënat mblidhen dhe përpunohen nga sensorët fizikë, ato transferohen te Binjaku Dixhital, i cili ndërton një pasqyrim dinamik të gjendjes aktuale të ambientit. Ky model virtual shërben si bazë për analizën parashikuese të realizuar nga modelet e Inteligjencës Artificiale, duke mundësuar simulimin dhe vlerësimin e gjendjeve të ardhshme të lagështirës. Në këtë mënyrë, Binjaku Dixhital funksionon si një ndërmjetës inteligjent ndërmjet botës fizike dhe sistemit vendimmarrës, duke integruar monitorimin në kohë reale, analizën e të dhënave, parashikimin dhe kontrollin automatik në një mekanizëm të vetëm funksional. Përmes mekanizmit të reagimit, rezultatet e veprimit fizik rikthehen sërish në sistem, duke i mundësuar Binjakut Dixhital të përditësojë vazhdimisht gjendjen virtuale dhe të mbështesë vetë-përshtatjen dinamike të sistemit.

6.3. Sinkronizimi me sensorët real

Shtresa e mbledhjes së të dhënave është komponenti fillestar dhe shumë i rëndësishëm i arkitekturës që përfshin DT, duke mundësuar lidhjen direkte midis sistemit fizik dhe përfaqësimit të tij virtual. Kjo shtresë është përgjegjëse për kapjen, trasmetimin dhe ruajtjen fillestare të të dhënave të gjeneruara nga një pajisje IoT për monitorimin e cilësisë së ajrit në kohë reale. Qëllimi kryesor i kësaj shtrese është të sigurojë një rrjedhë

të vazhdueshme dhe të besueshme të të dhënave bruto, pa ndërhyrje në strukturën apo përmbajtjen e tyre, duke garantuar kështu integritetin e informacionit që do të përdoret në fazat pasuese të përpunimit dhe analizës.

Pajisja e përdorur që përfaqëson një sistem të avancuar IoT për monitorimin ambjental është IQAir AirVisual Indoor/Outdoor (Figura 6.2) e projektuar për matjen e disa prej parametrave kryesor të cilësisë së ajrit dhe trasmetimin e tyre në kohë reale drejt një platforme cloud.



Figura 6.2 Pajisja AirVisual dhe rezultatet e aplikacionit online

Pajisja përfshin komponentë të ndryshëm fizikë, përfshirë sensorë të dedikuar dhe module të ndërrueshme për matjen e grimcave të ndotjes. Adresat MAC (Media Access Control Address) mundësojnë komunikimin në rrjet për lidhjen me Ethernet, WiFi dhe USB4G, ndërkohë me anë të indikatorëve led mundësohet informimi koherent mbi statusin e funksionimit të pajisjes. Të gjitha këto karakteristika e bëjnë pajisjen të përshtatshme për integrim në arkitektura të tipit Sistem Kibernetiko-Fizik dhe DT. Shtresa e mbledhjes së të dhënave luan një rol kritik në funksionimin e DT, pasi ajo përfaqëson pikën e kontaktit midis botës fizike dhe asaj digjitale. Pa këtë shtresë, nuk do të ishte e mundur sinkronizimi i gjendjes reale të sistemit me modelin virtual. Në këtë kontekst, pajisja e përdorur shërben si:

- burimi primar i të dhënave
- ndërfaqja midis sistemit fizik dhe modelit virtual
- input për modelet e inteligjencës artificiale dhe mekanizmat e vendimmarrjes

Parametrat që vlerësohen nga pajisja janë grimcat,

- PM2.5 ($\mu\text{g}/\text{m}^3$) që janë ≤ 2.5 mikrometër dhe zakonisht burimi i tyre është djegia e karburantit, ndotja urbane, automjetet etj. Saktësia është $\pm 10 \mu\text{g}/\text{m}^3$ ose $\pm 10\%$, dhe është ndër grimcat që lidhet drejtpërdrejtë me sëmundje respiratore dhe variabli më i rëndësishëm për Indeksin e Cilësisë së Ajrit (AQI)
- PM10 ($\mu\text{g}/\text{m}^3$) janë grimca $\leq 10 \mu\text{m}$ dhe burimi i tyre është pluhuri rrugës, ndërtimet dhe aktivitetet industriale. Këto grimca ndikojnë në rrugët e sipërme të frymarrjes dhe klasifikohen si më pak të rrezikshme se PM 2.5.
- PM1.0 ($\mu\text{g}/\text{m}^3$) Grimca shumë të imëta me diametër ≤ 1 mikrometër (μm) karakteristike të tyre është se depërtojnë thellë në mushkëri dhe mund të hynë në qarkullimin e gjakut. Ato kategorizohen si shumë të rrezikshme për shëndetin dhe përdoren në analiza të avancuara të ndotjes.
- Lagështira që tregon përqindjen e avullit të ujit në ajër dhe ruhet në (%). Saktësia është saktësia $\pm 1\%$. Ky parametër është shumë kritik për shëndetin, mykun dhe cilësinë e ajrit. Gjatë këtij rasti studimor kë është marrë si variabli kryesor në vendimmarrje.

Intervali i vlerave të këtij parametri është 0 – 100% Lagështirë Relative (pra sa i mbushur është ajri me avull uji krahasuar me maksimumin që mund të mbajë).

- Presionin atmosferik që tregon presionin e ajrit në atmosferë në hectopascal (hPa). Sensori mund të masë presionin e ajrit nga 300 hPa deri në 1100 hPa. Vlera 1013 hPa konsiderohet si presion normal, ndërkohë vlera <1000 hPa është presion i ulët, vlera >1020 hPa konsiderohet presion i lartë.
- Temperatura një parametër që përfshin intervalin -40°C deri në 90°C dhe me saktësi $\pm 2^\circ\text{C}$. Temperatura ndikon në lagështirë, në përhapjen e ndotësve si dhe në komfortin e njeriut.

Ndërkohë për të bërë një paralelizëm u ndërtua një pajisje Arduino me sensorin e PM 2.5 të lidhur veçmas e cila vlerësonte lokalisht nivelet e kësaj grimce. Mjeti i integruar Arduino për matjen e cilësisë së ajrit, i përdorur në këtë studim, përfshin vlerësimin e parametrave PM2.5, PM1.0, PM10 dhe AQI (Indeksi i Cilësisë së Ajrit).

Elementët përbërës të këtij sistemi të integruar janë si më poshtë:

- Grove - Laser PM2.5 Dust Sensor - Arduino Compatible - HM3301, i cili është një sensor për matjen e nivelit të grimcave PM2.5, të klasifikuara si të rrezikshme për shëndetin e njeriut kur tejkalojnë kufijtë normalë.
- Grove Base Shield V2.0 për Arduino, i cili shërben si bazë lidhëse ndërmjet pajisjes Arduino Uno dhe versioneve të tjera si Arduino Leonardo dhe Arduino Mega. Lidhësit me 4 pina mundësojnë lidhje të thjeshta dhe të shpejta.
- Seeeduno V4.2 (ATMega328P), që përfaqëson pllakën Arduino të pajisur me mjedisin e zhvillimit për programim dhe kontroll të sistemit.
- Grove - Universal 4 Pin Buckled 20 cm Cable, i cili përdoret për lidhjen ndërmjet Base Shield, Arduino dhe sensorit PM2.5.
- Lipo Rider Plus (Charger/Booster) - 5V/2.4A USB Type-C, që shërben si komponent lidhës ndërmjet pllakës Arduino dhe baterisë litium, duke mundësuar furnizimin me energji të sistemit.

Marrëdhënia ndërmjet këtyre komponentëve dhe mënyra e integritit të tyre paraqitet në figurën e mëposhtme (Figura 6.3).

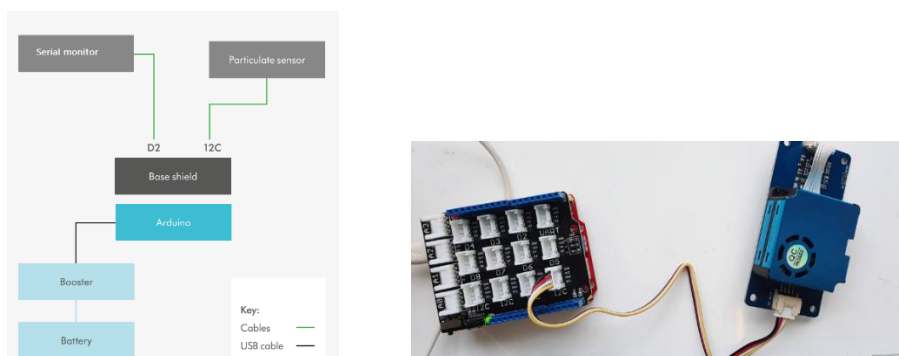


Figura 6.3 Arkitektura e pajisjes Arduino dhe modelimi fizik

Fillimisht matjet u realizuan në disa ambjente të Fakultetit të Shkencave të Natyrës më shumë për të paralelizuar rezultatet e dy pajisjeve si dhe vlerësimin e cilësisë së ajrit nga vetë studentët. Me anë të një pyetësi online u identifikua gjendja shëndetësore momentale dhe të përgjithshme e studentëve me anë të perceptimit, pra asaj që ato ndjenin në ajrin e klasave apo laboratoreve, si dhe situatës së tyre shëndetësore. Të gjitha këto veprimtari u skematizuan sipas figurës 6.4 ku matjet me pajisjet dhe plotësimi i pyetësorit online u realizuan paralelisht.

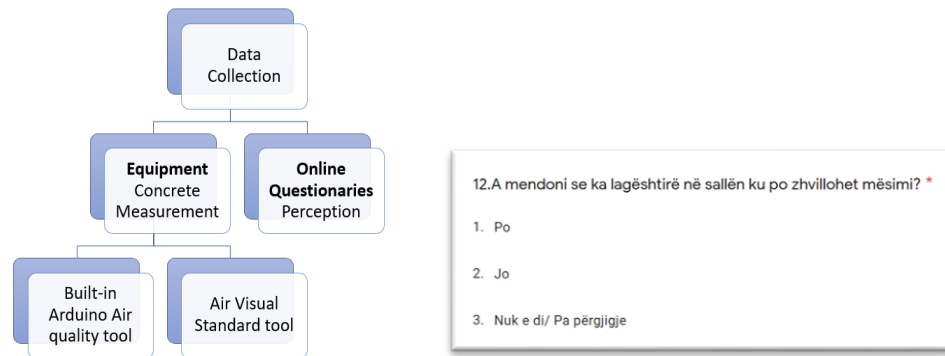


Figura 6.4 Struktura e veprimtarisë dhe pyetësori online

Problematika kryesore qëndronte tek mospërputhja e rezultateve të pajisjes Arduino me Pajisjen IQAir me një diferencë të konsiderueshme. Problematikë tjetër ishte edhe mungesa e disa pajisjeve të tjera me sensorët e matjes që mund të vendoseshin paralelisht në ambjente të ndryshme. Duke marrë shkas nga fakti që IQAir është një pajisje e bazuar në standartet si FCC (Federal Communications Commission-USA) i cili garanton që pajisja përmbush kufijtë e emetimeve elektromagnetike dhe mund të përdoret në mënyrë të sigurt në rrjetet e komunikimi. UL (underwriters Laboratories) është standard që rrit besueshmërinë sidomos në rrezikun elektrik. Standarti RoHS (Restriction of Hazardous Substances) si pajisje që nuk përmban substanca të rrezikshme. Standarti CE (Conformité Européenne) ku përmbushen standartet e Bashkimit Europian për sigurinë, shëndetin dhe mjedisin, dhe pa këtë standard pajisja nuk mund të shitet në BE. Kombinimi i të gjithë këtyre standarteve siguron cilësi të mirë të sinjalit, reduktim të gabimeve hardware, stabilitet në matje dhe trasmetim korrekt të të dhënave. Si rezultat i këtyre diferencave u vendos të përdorej IQAir si pajisje që ofron rezultate më të sakta.

Pajisja IQAir transmeton të dhënat në një platformë cloud të dedikuar. Kjo platformë shërben si një ndërmjetës midis pajisjes fizike dhe sistemit të DT (Digital Twin), duke mundësuar ruajtjen, përpunimin dhe aksesimin e të dhënave në kohë reale. Edhe pse në implementimin praktik të sistemit të DT, të dhënat aksesohen përmes çelësi (API key), platforma cloud mbetet komponent kyç në arkitekturën e sistemit, pasi ajo menaxhon komunikimin me pajisjen fizike dhe siguron integritetin dhe disponueshmërinë e të dhënave. Si konkluzion ky rast studimor përshtat një arkitekturë indirekte të marrjes së të dhënave, ku sensorët nuk komunikojnë drejtpërdrejtë me sistemin DT, por përmes një platforme cloud që ofron API për integrim (Figura 6.5). Të dhënat historike të disponueshme online janë përdorur kryesisht për analizimin e strukturës së të dhënave, identifikimin e problematikave, kuptimin e trendit kohor dhe vlerësimin e marrëdhënieve

ndërmjet parametrave mjedisorë. Ky historik i zgjeruar mundësoi analizën statistikore, identifikimin e sezonalitetit dhe vlerësimin e sjelljes dinamike të sistemit. Megjithatë, në implementimin praktik të sistemit të vendimmarrjes automatike dhe DT, parashikimi i lagështirës për horizontin 30-minutësh nuk kërkon domosdoshmërisht një histori shumë të gjatë të të dhënave, pasi modelet bazohen kryesisht në sjelljen afatshkurtër dhe varësitë kohore të intervaleve më të afërta.

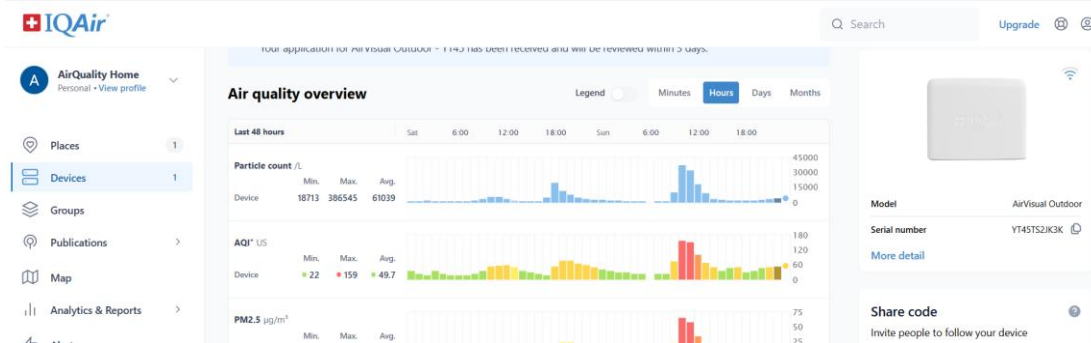


Figura 6.5 Platforma në cloud e IQAir

6.4. Shtresa e përpunimit dhe ruajtjes së të dhënave

Shtresa e përpunimit dhe ruajtjes së të dhënave përfaqëson komponentin ndërmjetës ndërmjet shtresës fizike të monitorimit dhe komponentëve inteligjentë të sistemit DT. Qëllimi kryesor i kësaj shtrese është marrja, organizimi, përpunimi dhe ruajtja e të dhënave të gjeneruara nga sensorët mjedisorë, duke siguruar që informacioni i përdorur nga sistemi të jetë i besueshëm, i sinkronizuar dhe i përshtatshëm për analizë të mëtejshme.

Në sistemin e propozuar, të dhënat përditësohen periodikisht çdo 5 minuta dhe ruhen në formë të strukturuar së bashku me informacionin kohor përkatës. Duke qenë se të dhënat burojnë nga pajisje IoT dhe procese reale monitorimi, ato mund të përmbajnë mungesa, ndërprerje komunikimi, vonesa kohore ose vlera jo të rregullta të lidhura me kushtet e funksionimit të sensorëve dhe rrjetit të komunikimit.

Në vijim të analizës së strukturës së të dhënave të paraqitur në Kapitullin III, dataset-i u konsiderua i përshtatshëm për analizë përpara integritetit në sistemin Digital Twin, pasi paraqet frekuencë të rregullt matjeje, nivel të ulët mungesash, përqindje të pranueshme anomalish dhe varësi të qartë kohore ndërmjet matjeve. Megjithatë, përpara ruajtjes në SQLite u aplikuan transformime jo shumë të thella modeluese, si krijimi i veçorive me vonesë kohore, normalizimi apo ndërtimi i target-it parashikues, sepse këto hapa lidhen drejtpërdrejt me fazën e trajnimit dhe përdorimit të modeleve të Inteligjencës Artificiale. Në fazën e ruajtjes u kryen vetëm procese bazë të standardizimit dhe kontrollit të cilësisë, si unifikimi i emrave të kolonave, konvertimi i timestamp-it në format të qëndrueshëm, sigurimi i tipit numerik për temperaturën, lagështirën dhe trysninë, si dhe shmangia e rritjes së pakontrolluar të dataset-it përmes mekanizmit “sliding window”. Në këtë mënyrë, SQLite shërben si shtresë lokale e ruajtjes dinamike të të dhënave të DT, ndërsa përpunimi analitik dhe krijimi i veçorive parashikuese realizohen në një fazë të mëvonshme, kur të dhënat lexohen nga databaza dhe përdoren si input për modelet parashikuese dhe modulën e vendimmarrjes automatike.

Duke marrë në konsideratë faktin që Binjaku Dixhital duhet të ketë të dhëna në kohë reale u implementua në gjuhën Python një script automatik dhe u integrua me API-në e platformës IQAir/AirVisual për të marrë parametrat aktualë atmosferikë, si temperatura, lagështira dhe trysnia atmosferike. Skripti ekzekutohet automatikisht çdo 5 minuta përmes Windows Task Scheduler, Figura 6.6, duke simuluar një proces të vazhdueshëm sinkronizimi midis mjedisit fizik dhe përfaqësimit virtual të tij në DT.

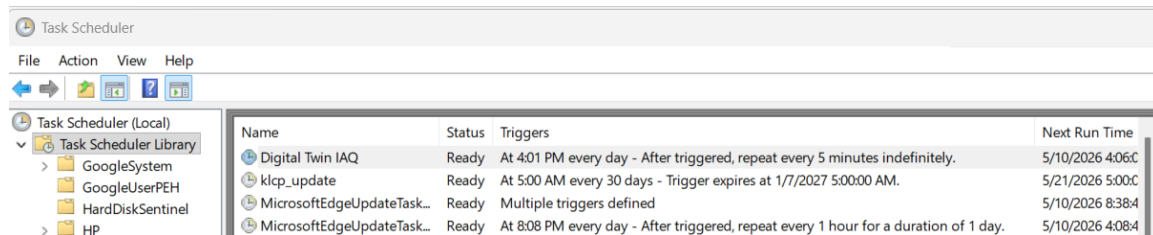


Figura 6.6 Krijimi i skedulimit për ekzekutimin e skriptit në Task Scheduler të Windows

Në çdo cikël ekzekutimi, sistemi lexon dataset-in ekzistues, shton rekordin më të ri të marrë nga API në fund të dataset-it dhe fshin automatikisht rekordin më të vjetër në fillim të tij (mekanizmi Sliding Window). Për të shmangur rritjen e pakontrolluar të madhësisë së dataset-it dhe për të optimizuar performancën e sistemit, u vendos një kufi maksimal prej 10,000 rreshtash, afërsisht i një periudhe 1 muaj, Figura 6.7 .

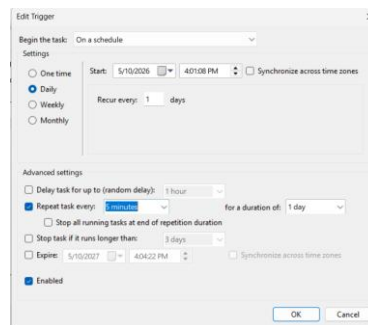


Figura 6.7 Trigeri i detyrës automatik i konfiguruar çdo 5 min

Fillimisht, mekanizmi i përditësimit të të dhënave u realizua duke përdorur një dataset në format CSV. Në çdo cikël ekzekutimi, skripti shtonte matjen më të fundit në fund të skedarit dhe largonte rekordin më të vjetër, duke zbatuar logjikën e një dritareje rrëshqitëse. Dataset-i i përditësuar më pas përdorej si input për modelet parashikuese dhe modulën e vendimmarrjes automatike.

Megjithatë, përdorimi i CSV-së paraqiti kufizime praktike, pasi çdo përditësim kërkonte leximin, modifikimin dhe rishkrimin e të gjithë skedarit. Kjo qasje bëhet më pak efikase kur të dhënat përditësohen periodikisht dhe kur sistemi duhet të funksionojë në mënyrë të vazhdueshme. Për këtë arsye, u kalua në përdorimin e databazës SQLite, e cila mundëson ruajtje më të qëndrueshme, përditësim më efikas të rekordeve dhe menaxhim më të mirë të mekanizmit të dritares rrëshqitëse.

Në SQLite, matja e re shtohet si rekord i ri në tabelë, ndërsa rekordi më i vjetër fshihet automatikisht kur tejkalohet numri maksimal i rreshtave të përcaktuar. Në këtë mënyrë, databaza ruan vazhdimisht vetëm historikun më të fundit të nevojshëm për

parashikimin 30-minutësh dhe vendimmarrjen automatike në kuadër të DT, Figura 6.8.

	id	Datetime	Temperatura	Lageshtira	Trysnia	collected_at
11	12	12/1/2025 0:25	17.1	80.6	100067.0	2026-05-10 18:37:28.186580
12	13	12/1/2025 0:30	17.2	81.0	100064.6	2026-05-10 18:37:28.186580
13	14	12/1/2025 0:35	17.2	81.0	100064.2	2026-05-10 18:37:28.186580
14	15	12/1/2025 0:40	17.3	81.0	100064.2	2026-05-10 18:37:28.186580
15	16	12/1/2025 0:45	17.3	81.0	100065.0	2026-05-10 18:37:28.186580
16	17	12/1/2025 0:50	17.3	81.0	100066.0	2026-05-10 18:37:28.186580
17	18	12/1/2025 0:55	17.3	81.0	100066.8	2026-05-10 18:37:28.186580
18	19	12/1/2025 1:00	17.3	81.0	100071.0	2026-05-10 18:37:28.186580
19	20	12/1/2025 1:05	17.3	81.0	100078.6	2026-05-10 18:37:28.186580

Figura 6.8 Databaza SQLite

E njëjta logjikë që u përdor për ruajtjen e të dhënave në mënyrë koherente në formatin csv është përdorur edhe për skedarin db të SQLite. Implementimi është realizuar në python. Për të realizuar komunikimin me databazën lokale, marrjen e të dhënave nga API dhe përpunimin e tyre u përdorën libraritë sqlite3, requests dhe pandas. Skedari digital_twin.db shërben për burimin kryesor të të dhënave të Digital Twin. Skripti që është skeduluar për t'u ekzekutuar çdo 5 min përfshin disa funksionalitete. Fillimisht krijohet tabela nëse nuk ekziston, ndërsa funksioni insert_record() krijon një lidhje me API e pajisjes IQAir dhe merr të dhënat e fundit të lagështirës, temperaturës dhe trysnisë. Të dhënat e reja shtohen në databazë përmes komandës SQL INSERT INTO, duke simuluar sinkronizimin periodik midis mjedisit fizik dhe modelit virtual. Funksioni delete_first_row() fshin automatikisht rekordin më të vjetër në databazë pas shtimit të rekordit të ri. Funksioni show_last_rows() përdoret për monitorim dhe shfaq pesë rekordet më të fundit të ruajtura në sistem. Në fund, funksioni run_script() orkestron të gjithë procesin, duke realizuar krijimin e tabelës, marrjen e të dhënave, ruajtjen në databazë dhe fshirjen e rekordit më të vjetër në një cikël të vetëm ekzekutimi, i cili mund të planifikohet periodikisht përmes Windows Task Scheduler i cili realizon përditësimin automatik çdo 5 minuta. Figura 6.9 paraqet mesazhin informues mbi ndryshimin e bazës së të dhënave.

```

C:\Program Files\WindowsAp...
Po ekzekutohet Digital Twin me SQLite...
U shtua rekordi i ri në SQLite.
U fshi rreshti i parë më i vjetër.
5 rreshtat e fundit:
id      Datetime      Temperatura  Lageshtira  Trysnia      collected_at
0 8915 2026-05-10 16:00:00      22.0        59.0        1014.0      2026-05-10 19:07:10.652950
1 8914 2026-05-10 16:00:00      22.0        59.0        1014.0      2026-05-10 18:45:33.137003
2 8913 2026-05-10 16:00:00      22.0        59.0        1014.0      2026-05-10 18:42:56.798981
3 8912 2026-05-10 16:00:00      22.0        59.0        1014.0      2026-05-10 18:42:19.497378
4 8911 2026-05-10 16:00:00      22.0        59.0        1014.0      2026-05-10 18:40:42.081541
U krye me sukses.
Shtyp ENTER për të mbyllur dritaren...

```

Figura 6.9 Skedulimi taskut ne SQLite

6.5. Integrimi modeleve të Inteligjencës Artificiale

Shtresa e modelit parashikues ka si qëllim të parashikojë vlerën e ardhshme të lagështirës në ambient të brendshëm duke përdorur të dhënat historike të mbledhura nga sensori. Rezultati i kësaj shtrese shërben si input për modulën e vendimarrjes automatike, i cili aktivizon ose çaktivizon pajisjen e kontrollit të lagështirës në varësi të pragjeve të përcaktuara. U testuan disa teknika aritmetike, ML (Machine Learning) dhe DL (Deep Learning) për parashikimin dhe përqindjen më të lartë në saktësi e cila u vendos edhe si kryesore në lidhje me modulën e vendimarrjes. Përpara se të aplikohen teknikat u përcaktua se si do gjendej intervali kohor parashikues më optimal si dhe gjetja e pragut kritik për aktivizimin e vendimit.

6.5.1. Gjetja e intervalit kohor parashikues

Për të përcaktuar frekuencën reale të kampionimit të të dhënave kohore, u analizuan diferencat midis stampave kohore të njëpasnjëshme. Konkretisht, për çdo çift vlerash kohore t_i dhe t_{i-1} , u llogarit diferenca:

$$\Delta t_i = t_i - t_{i-1} \quad (6.1)$$

Më pas, për të shmangur ndikimin e vlerave ekstreme, u përdor mediana e këtyre diferencave si një vlerësim robust i intervalit të kampionimit. Ky interval u konvertua në minuta dhe u përdor për të përcaktuar frekuencën dominante të të dhënave.

Përcaktimi i saktë i këtij intervali është thelbësor për ndërtimin e modeleve të parashikimit në seri kohore, pasi ndikon drejtpërdrejt në zgjedhjen e horizontit të parashikimit dhe strukturën e variablave të vonuar. Për të përshtatur horizontin e parashikimit me strukturën diskrete të të dhënave kohore, u realizua konvertimi nga njësi kohore (minuta) në hapa diskretë (hapa). Ky konvertim u bazua në intervalin e kampionimit të përcaktuar më parë. Konkretisht, numri i hapave për një horizont të dhënë H (në minuta) u llogarit sipas formulës:

$$\text{hapat} = \frac{H}{\Delta t} \quad (6.2)$$

ku Δt përfaqëson intervalin e kampionimit në minuta. Për të siguruar që rezultati të jetë një numër i plotë, u aplikua rrumbullakosja në vlerën më të afërt:

$$\text{hapat} = \text{round}\left(\frac{H}{\Delta t}\right) \quad (6.3)$$

Ky transformim mundëson përdorimin e horizonteve të ndryshëm të parashikimit (p.sh. 15, 30 dhe 60 minuta) në modele të serive kohore, të cilat operojnë mbi indekse diskrete kohore. Gjithashtu, u aplikua një kontroll vlefshmërie për të siguruar që numri i hapave të jetë më i madh se zero, duke garantuar konsistencën e të dhënave dhe korrektësinë e intervalit të kampionimit.

Për të përcaktuar horizontin optimal të parashikimit për lagështirën, u analizua ndryshimi absolut i vlerave në intervale të ndryshme kohore (15, 30 dhe 60 minuta).

Ndryshimi u llogarit sipas formulës:

$$| H(t + \Delta t) - H(t) | \quad (6.4)$$

ku $H(t)$ përfaqëson lagështirën në kohën aktuale dhe $H(t + \Delta t)$ vlerën në të ardhmen. Tabela 6.1 paraqet procesin e analizës së variacionit të lagështirës për tre horizonte të ndryshme parashikues 15,30 dhe 60 minuta. Duke u bazuar në të dhënat ekzistuese të lagështirës të regjistruara në intervale kohore periodike, u krijuan kolonat e reja dif_15, dif_30 dhe dif_60, të cilat përfaqësojnë ndryshimin absolut të lagështirës ndërmjet vlerës aktuale dhe vlerës së ardhshme pas secilit horizont kohor. Rezultatet tregojnë se ndryshimet në horizontin 15-minutësh janë relativisht të vogla, duke reflektuar stabilitet afatshkurtër të ambientit, ndërsa horizontet 30 dhe 60 minuta shfaqin variacione më të dukshme. Kjo analizë shërben si bazë për vlerësimin e sjelljes dinamike të sistemit dhe për identifikimin e horizontit më të përshtatshëm të parashikimit në kuadër të modelit DT për ndihmesën e realizimit të vendimmarrjes automatike në menaxhimin e lagështirës.

Tabela 6.1 Analiza e variacionit të tre horizonteve parashikues

Kolonat ekzistuese		Kolonat e reja		
Data - Ora	Lagështira	dif_15	dif_30	dif_60
12/1/2025 0:00	80.0	0.0	1.0	1.0
12/1/2025 0:05	80.0	0.2	1.0	1.0
12/1/2025 0:10	80.0	0.6	1.0	1.0
12/1/2025 0:15	80.0	1.0	1.0	1.0
12/1/2025 0:20	80.2	0.8	0.8	0.8
12/1/2025 0:25	80.6	0.4	0.4	0.4

Rezultatet statistikore të variacioneve të lagështirës për horizontet parashikuese 15, 30 dhe 60 minuta paraqiten në Tabelën 6.2. Për secilin horizont janë llogaritur mesatarja, mediana, devijimi standard, varianca, si dhe vlerat minimale dhe maksimale të ndryshimeve të lagështirës. Rezultatet tregojnë se me rritjen e horizontit kohor rritet edhe niveli i variacionit të të dhënave. Horizonti 15-minutësh paraqet ndryshime më të vogla dhe stabilitet më të lartë, ndërsa horizonti 60-minutësh shfaq devijim standard dhe variancë më të lartë, duke reflektuar rritjen e pasigurisë dhe paqëndrueshmërisë së sistemit në parashikime afatgjata. Horizonti 30-minutësh paraqet një balancë ndërmjet stabilitetit dhe aftësisë për të kapur ndryshime të rëndësishme të lagështirës, çka e bën atë një alternativë të përshtatshme për modelimin parashikues dhe ndihmesë në realizimin e vendimmarrjes automatike në kuadër të sistemit DT.

Tabela 6.2 Rezultatet statistikore të horizonteve parashikuese

Horizonti	Mesatarja	Mediana	Devijacioni standard	Varianca	Min	Max
15 min	0.337508	0.0	0.898329	0.806996	0.0	24.0
30 min	0.522555	0.0	1.160947	3.380911	0.0	23.6

60 min	0.791664	0.4	1.362284	5.923547	0.0	23.2
--------	----------	-----	----------	----------	-----	------

Për të bërë zgjedhjen finale të intervalit kohor parashikues duhet konsideruar edhe fakti që sa mirë parashikon modeli për secilin horizont? Kjo u realizua duke nxjerrë në pah vlerat e metrikave si RMSE (Root Mean Squared Error) . Ajo mat ndryshimin e parashikimet nga vlera reale, duke penalizuar më shumë gabimet e mëdha sepse i ngre në katror , për shembull rasti i një gabimi me diferencë 10 ajo e quan si 100 gjë që është ndryshim shumë i madh. Formula e RMSE është si mëposhtë:

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (6.5)$$

Ku y_i = vlera reale , $\hat{y}_i$ = vlera e parashikuar , n = numri i mostrave.

Rastet kur kemi RMSE të vogla konkludohet që modeli është shumë i mirë. Modeli RMSE penalizon shumë gabimet e mëdha.

Metrika tjetër është MAE (Mean Absolute Error) e cila është mesatarja e gabimeve absolute (pa i ngritur në katror) dhe kuptimi i vlerave ngjason me RMSE

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad (6.6)$$

Interpretimi i kësaj metrike është që vlera sa më të vogla të MAE nënkuptojnë që aq më i mirë do jetë modeli. Në ndryshim nga RMSE , MAE nuk penalizon shumë gabimet e mëdha.

Metrika e tretë është R^2 (R-squared-Coefficient of Determination) i cili mat se sa mirë modeli shpjegon variancën e të dhënave . Formula e metrikës R^2 është:

$$R^2 = 1 - \frac{\sum (y_i - \hat{y}_i)^2}{\sum (y_i - \bar{y})^2} \quad (6.7)$$

Ku $\bar{y}$ është mesatarja e vlerave reale dhe interpretimi i rezultateve të metrikës është:

$R^2 = 1 \rightarrow$ model perfekt

$R^2 = 0 \rightarrow$ model si mesatarja

$R^2 < 0 \rightarrow$ model shumë i dobët

Për secilin interval u aplikuan modele regressive dhe ML (Machine Learning) dhe u realizua krahasimi i këtyre metrikave.

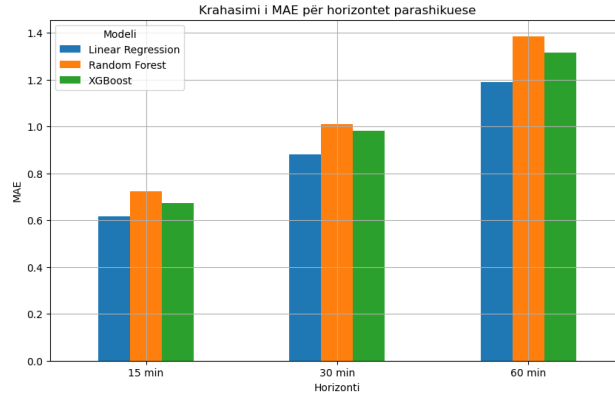


Figura 6.10 Metrika MAE për horizonte nga tre modelet

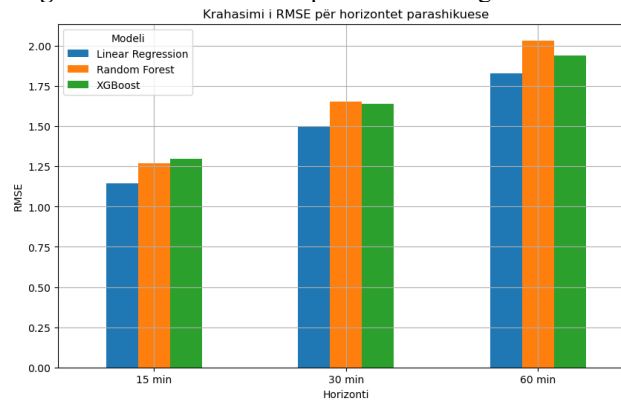


Figura 6.11 Metrika RMSE për horizonte nga tre modelet parashikuese

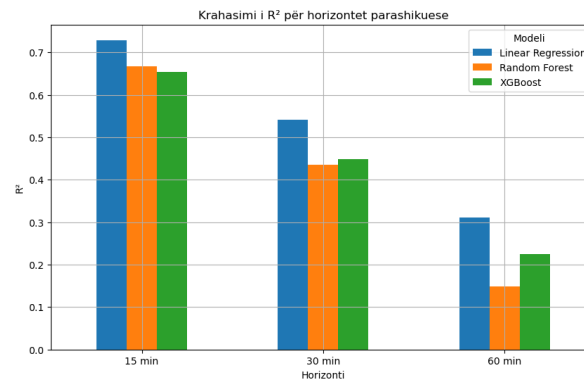


Figura 6.12 Metrika R^2 për horizonte nga tre modelet parashikuese

Analiza krahasuese e metrikave (Figura 6.12) MAE, RMSE dhe R^2 për tre horizonte parashikuese tregon se performanca e modeleve bie gradualisht me rritjen e horizontit kohor. Horizonti 15-minutësh paraqet gabimet më të ulëta dhe vlerat më të larta të R^2 , duke reflektuar stabilitet dhe parashikueshmëri të lartë afatshkurtër. Megjithatë, ky horizont ofron kohë të kufizuar për realizimin e vendimmarrjes proaktive. Në anën tjetër, horizonti 60-minutësh shfaq rritje të konsiderueshme të gabimeve dhe ulje të performancës së modeleve, duke treguar pasiguri më të lartë në parashikime afatgjata. Horizonti 30-minutësh përfaqëson kompromisin më të përshtatshëm ndërmjet saktësisë së modelit dhe aftësisë së sistemit për reagim proaktiv, duke u konsideruar si alternativa optimale për sistemin Digital Twin të menaxhimit automatik të lagështirës.

Bazuar në analizën e frekuencës së kampionimit, variacionit statistikor të lagështirës dhe performancës së modeleve parashikuese përmes metrikave MAE, RMSE dhe R^2 , arrihet në përfundimin se horizonti 30-minutësh përfaqëson alternativën më të përshtatshme për sistemin DT të menaxhimit automatik të lagështirës. Edhe pse horizonti 15-minutësh paraqet gabimet më të ulëta dhe stabilitet më të lartë, ai ofron kohë të kufizuar për realizimin e kontrollit proaktiv. Në anën tjetër, horizonti 60-minutësh shfaq rritje të ndjeshme të pasigurisë dhe ulje të performancës së modeleve. Horizonti 30-minutësh arrin një balancë ndërmjet saktësisë së parashikimit, stabilitetit të sistemit dhe kohës së nevojshme për reagimin automatik të pajisjes, duke u konsideruar si horizonti optimal për vendimmarrjen automatike në kuadër të Binjakut Dixhital.

6.5.2. Përcaktimi i pragut kritik të lagështirës relative

Zgjedhja e pragut për lagështirën u realizua duke kombinuar analizën statistikore të të dhënave me kuptimin praktik të sistemit të menaxhimit të ambientit të brendshëm. Fillimisht u analizua shpërndarja e vlerave të lagështirës përmes statistikave përshtatshme dhe percentileve (pozicioni relativ të një vlere brenda një shpërndarjes së të dhënave), me qëllim identifikimin e intervaleve ku ndodhet shumica e vlerave. Më pas, u testuan disa pragje të mundshme (p.sh. 55%, 60%, 65%) për të vlerësuar ndikimin e tyre në ndarjen e të dhënave në klasa (mbi dhe nën prag). Zgjedhja përfundimtare e pragut u bazua në një kompromis midis kuptimit fizik (niveli i lagështirës ku kërkohet ndërhyrje, p.sh. aktivizimi i thithësit të lagështirës) dhe shpërndarjes së balancuar të klasave, duke siguruar që modeli parashikues të jetë njëkohësisht i saktë dhe i dobishëm.

Duke marrë në konsideratë ASHRAE (American Society of Heating, Refrigerating and Air-Conditioning Engineers) Standard 55-2017 (Thermal Environmental Conditions for Human Occupancy) që është më i përdorur në sistemet HVAC, Smart Buildings etj sugjeron që lagështira zakonisht duhet të jetë $\leq 60\%$ dhe si zonë komforti janë vlerat rreth 30–60% [172]. Intervale të ngjashme trasmetohen edhe nga Standardi EN 16798 (Indoor Environmental Quality) që gjendje komforti optimale të lagështirës është rangi 40%-60% [173]. Vlerat e lagështirës kryesisht në rangun 60%-75 % ndikojnë në shtimin e mykut dhe ky i fundit ndikon negativisht në rrugët e frymarrjes duke shkaktuar alergji dhe reaksione inflamatorë [174].

Tabela 6.3 paraqet shpërndarjen e lagështirës relative sipas muajve të vitit 2025 dhe tregon se sjellja e sistemit ndryshon ndjeshëm gjatë vitit, duke reflektuar praninë e sezonalitetit në të dhënat e monitoruara dhe regjistruara. Vlerat e pragut të balancuar ndryshojnë nga rreth 43% gjatë muajve të verës deri në mbi 80% në periudhën dimërore, çka tregon se një prag unik nuk përfaqëson në mënyrë optimale të gjitha kushtet klimatike. Këto rezultate sugjerojnë mundësinë e përdorimit të pragjeve me përshtatje mujore në sisteme më të avancuara të vendimmarrjes automatike dhe DT. Megjithatë, për qëllime të stabilitetit operacional, interpretueshmërisë dhe përputhshmërisë me standardet HVAC (Heating, Ventilation and Air Conditioning) dhe Cilësinë e ajrit në ambjentet e brendshme, në këtë studim u përshtat një prag fiks fillimisht prej 69% për aktivizimin e mekanizmit të kontrollit të lagështirës.

Tabela 6.3 Shpërndarja e lagështirës sipas muajve të vitit 2025

Muaji	Mesatarja	Mediana	Pragu_balancuar	Nen_prag	Mbi_prag
1	78.59	80	80.2	2533	2353
2	77.63	80	80.2	4049	4006
3	74.16	76	76.2	4636	4196
4	70.8	73	73.2	4518	4080
5	62.36	63	63.2	3757	3680
6	49.22	50.2	50.4	4323	4306
7	44.36	43	43.2	4468	4460
8	48.87	48	48	4173	4331
9	62.68	64	64.2	4340	4042
10	68.05	69	69	4370	4557
11	75.12	76	76	4111	4529
12	76.66	77	77	4184	4719

6.5.3. Aplikimi modeleve statistikore, të mësuarit të makinës dhe të mësuarit të thelluar

Analiza e strukturës së të dhënave e realizuar në Kapitullin III tregoi se dataset-i i lagështirës relative paraqet karakteristika tipike të serive kohore, duke përfshirë autokorrelacion, sezonalitet dhe varësi kohore ndërmjet matjeve të njëpasnjëshme. Për këtë arsye, përzgjedhja e teknikave të modelimit nuk u realizua në mënyrë rastësore, por u bazua në vetitë statistikore dhe strukturore të të dhënave.

Fillimisht u përfshi model statistikor dhe linear, si LR (Linear Regression), për të krijuar modele bazë interpretuese dhe për të vlerësuar sjelljen lineare të serisë kohore. Më pas u aplikuan modele ML (Machine Learning), si RF (Random Forest) dhe XGBoost, me qëllim kapjen e marrëdhënieve jolineare dhe ndërveprimeve komplekse ndërmjet variablave. Në vijim, u përdorën modele DL (Deep Learning), konkretisht LSTM dhe GRU, për shkak të aftësisë së tyre për të modeluar varësitë temporale afatshkurtra dhe afatgjata në seritë kohore.

Krahasimi i këtyre teknikave mundëson jo vetëm identifikimin e modelit me performancën më të lartë parashikuese, por edhe vlerësimin e kompromisit ndërmjet interpretueshmërisë dhe kompleksitetit për të arritur te realizimi i vendimarrjes automatike.

Duke qënë se qëllimi i modeleve është parashikimi i një vlere numerike të lagështirës dhe jo klasifikimi në kategori, vlerësimi i performancës së tyre bazohet në metrikat regresive si RMSE (gabimet e mëdha), MAE (gabimi mesatar), R^2 (sa mirë modeli shpjegon variacionin). Për testimin e modeleve parashikuese u përdorën biblioteka të ndryshme të ekosistemit Python, si pandas dhe numpy për përpunimin e të dhënave, scikit-learn për modelet bazë dhe metrikat e vlerësimit, xgboost për modelin XGBoost, tensorflow/keras për modelet LSTM dhe GRU, si dhe matplotlib për vizualizimin e rezultateve. Përpara aplikimit të modeleve u realizua krijimi i veçorive hyrëse, duke përfshirë vlerat aktuale të lagështirës, temperaturës dhe trysnisë, veçoritë me vonesë kohore, veçoritë sezonale dhe statistikave lëvizëse. Modelet tradicionale të ML (Machine Learning) kërkojnë që nxjerrja e veçorive si pjesë inxhinierike të jetë e qartë për të përfaqësuar varësitë kohore të serisë, ndërsa modelet DL (Deep Learning) përdorin

sekuenca kohore si input dhe mësojnë automatikisht marrëdhëniet kohore. Kjo qasje mundëson krahasimin objektiv të teknikave të ndryshme në të njëjtin problem parashikimi dhe vendimmarrjeje automatike.

Ideja kryesore e nxjerrjes së disa veçorive është qëllimi për ti dhënë modelit ‘kujtesë’ dhe ‘kontekst’. Pra modeli nuk sheh vetëm një numër por historinë dhe trendin. Veçoritë që janë konsideruar gjatë këtij studimi, për datasetin e lagështirës janë:

- Vlerat aktuale (të matjes)
- Veçoritë e vonesës (vlerat e mëparshme të lagështirës, sepse niveli lagështirës varet nga e kaluara dhe me anë të tyre sistemi kupton trendin)
- Veçoritë kohore (sepse lagështira ndryshon sipas orës, sezonit, ditës/natës nga e cila modeli mëson sjelle periodike)
- Veçoritë e rrotullimit (e cila ruan mesataren e lagështirës për 30 min e fundit-rasti studimor por edhe luhatshmërinë e lagështirës)
- Veçoritë e diferencës së lagështirës (është ndryshimi ose diferenca mes dy matjeve dhe me anë të saj modeli kupton nëse ka tendencë të rritet apo të zbresë lagështira)
- Veçoritë e diferencës së temperaturës (duke që se temperatura është në korrelacion të fortë me lagështirën)

Në rastin e Regresionit Linear, të Random Forest dhe XGBoost këto veçori duhet të krijohen, ndryshe për modelet LSTM dhe GRU ku modeli merr sekuencë kohore dhe mëson vetë trendin.

Për modelet parashikuese u realizua ndarja kronologjike e të dhënave në dataset trajnimi dhe dataset testimi. Dataset-i i trajnimit u përdor për mësimin e parametrave të modeleve, ndërsa dataset-i i testimit u përdor për vlerësimin objektiv të performancës mbi të dhëna të papara më parë. Në rastin e modeleve DL (Deep Learning), si LSTM dhe GRU, u përdor gjithashtu një dataset validimi për monitorimin e procesit të trajnimit dhe reduktimin e rrezikut të “mbipërshtatjes”. Duke qenë se problemi trajton seri kohore, ndarja e të dhënave u realizua në mënyrë kronologjike, për të shmangur “rrjedhjen e të dhënave” dhe për të ruajtur rendin kohor të matjeve. Në rastin e “rrjedhjes së të dhënave” modeli gjatë trajnimit merr informacion për të ardhmen që në realitet nuk duhet ta dijë, duke dhënë kështu rezultate artificialisht shumë të mira.

Dataset-i u nda në mënyrë kronologjike në dy pjesë: 80% e të dhënave u përdorën për trajnim dhe 20% për testim. Kjo ndarje ruan rendin temporal të serisë kohore dhe simulon skenarin real të përdorimit të sistemit, ku modeli trajnohet mbi të dhënat historike dhe përdoret për të parashikuar intervale kohore të ardhshme. Kjo qasje shmang fenomenin e “rrjedhjes së të dhënave” dhe siguron vlerësim më realist të performancës së modeleve parashikuese.

Duke qenë se lagështira relative paraqet sezonalitet të fortë, ndarja kronologjike train/test mund të krijojë ndryshime ndërmjet shpërndarjes së të dhënave të trajnimit dhe testimit, veçanërisht gjatë muajve dimërorë me nivele më të larta lagështire. Megjithatë, kjo qasje reflekton më mirë skenarin real të përdorimit të sistemit, ku modeli trajnohet mbi të dhëna historike dhe përdoret për të parashikuar periudha të ardhshme kohore. Për më tepër, rezultatet e analizës mujore sugjerojnë se pragjet statistikore të lagështirës ndryshojnë sipas sezonit, duke mbështetur perspektivën e përdorimit të pragjeve adaptive në zgjerimet e ardhshme të sistemit DT.

Për Regresionin Linear dhe Pylli i Rastësishëm numri i pemëve të vendimit është

përcaktuar 100 si një vlerë standarde që jep një stabilitet të mirë pa rritur kohën e ekzekutimit. Rëndësi ka edhe sasia e herëve që ekzekutohet kodi për të dhënë afërsisht të njëjtat rezultate dhe në rastin konkret është caktuar 42 (tipike për ML). Për realizimin e një trajnimit sa më të shpejtë është e mira të përdoren të gjithë procesorët CPU (Central Processing Unit) te `nr_jobs=-1`. Ndërkohë për XGBoost numri i pemëve është caktuar 150 sepse zakonisht kërkohet një sasi më e madhe se të Pylli i Rastësishëm dhe vlera 150 është një kompromis i mirë i performancës dhe kohës së trajnimit. Përsa i takon vlerës se sa shpejt mëson modeli (shkalla e të mësuarit) është vendosur 0.05. Për një shkallë të mësuarit të lartë modeli mëson shpejt por rrezikon mbipërshtatjen ndërsa një shkallë e vogël e të mësuarit bën që modeli të mësojë më ngadalë por në formë më të qëndrueshme.

Kompleksiteti i pemëve është caktuar 5 sepse nuk ka jo-linearitet ekstrem prandaj nuk nevojiten pemë shumë të thella.

Qëllimi kryesor i këtij studimi nuk është optimizimi i hiperparametrit për modelet për të arritur një performancë të lartë por është krahasimi i teknikave statistikore, të mësuarit të makinës dhe të mësuarit të thelluar në kuadër të sistemit DT (Digital Twin). Për këtë arsye, për secilin model u përdorën konfigurime fikse dhe të qëndrueshme të hiperparametrave, bazuar në vlera standarde të përdorura gjerësisht në literaturë dhe në eksperimente paraprake. Kjo qasje mundësoi vlerësimin objektiv të sjelljes së modeleve dhe krahasimin e tyre në kushte të njëjta eksperimentale, duke ruajtur interpretueshmërinë dhe kompleksitetin e kontrolluar të sistemit.

Modelet e të mësuarit të makinës dhe të mësuarit të thelluar u vlerësuan sipas dy pragjeve, fillimisht pragut fiks 69% dhe më pas pragu interval 58%–72%. Me anë të metrikave regressive u arritën disa rezultate që kishin ngjashmëri me njëra tjetrën. Duke mbajtur vlerën e horizontit 30 min metrikat MAE, RMSE dhe R^2 rezultuan me vlera sipas tabelave të mëposhtme. Tabela 6.4 paraqet vlerësimin e performancës së modeleve parashikuese për pragun fiks të lagështirës prej 69%. Rezultatet tregojnë se modeli i Regresionit Linear dhe modeli XGBoost kanë performancën më të mirë (Tabela 6.4) krahasuar me modelet e tjera, duke arritur vlera të ulëta të gabimit dhe një aftësi të lartë shpjeguese të variacionit të të dhënave.

Tabela 6.4 Vlerësimi i modeleve për pragun fiks 69%

Model	MAE	RMSE	R^2
Regresioni Linear	0.820897066	1.336579	0.896441
XGBoost	0.777922098	1.366821	0.891701
Pylli i rastësishëm	1.074898754	1.674496	0.837457
GRU	2.479291737	2.841445	0.532246
LSTM	2.865154806	3.167678	0.418672

Konkretisht, Regresioni Linear arrin vlerën më të lartë të $R^2 = 0.896$, çka tregon se modeli shpjegon rreth 89.6% të variacionit të lagështirës. XGBoost paraqet gjithashtu performancë shumë të afërt, me gabime minimale dhe stabilitet të mirë në parashikim. Ndërkohë, modeli Random Forest (Pylli i rastësishëm) ka performancë mesatare, me rritje të lehtë të gabimeve dhe ulje të koeficientit R^2 . Modelet e të mësuarit të thellë GRU dhe LSTM rezultojnë më pak efektive për këtë dataset dhe horizont parashikimi, duke shfaqur vlera më të larta të MAE dhe RMSE si dhe vlera më të ulëta të R^2 , gjë që tregon një aftësi më të

kufizuar në kapjen e dinamikës së lagështirës në krahasim me modelet klasike të ML. Në përgjithësi, rezultatet sugjerojnë se për pragun fiks 69%, modelet klasike të ML janë më të përshtatshme dhe më stabile sesa modelet DL për skenarin e analizuar.

Ndërsa rezultatet e paraqitura në Tabelën 6.5 paraqesin vlerësimin e të njëjtave metrika për të njëjtat modele por për pragun interval . Nga tabela kuptohet se modeli XGBoost arrin performancën më të balancuar në përgjithësi, duke regjistruar vlerën më të ulët të MAE (0.778), çka tregon gabimin mesatar më të vogël të parashikimit. Ndërkohë, Regresioni Linear paraqet vlerën më të ulët të RMSE (1.337) dhe vlerën më të lartë të R^2 (0.896), duke reflektuar aftësi shumë të mirë për të shpjeguar variacionin e të dhënave. Pylli i Rastësishëm paraqet performancë të moderuar, ndërsa modelet e të Mësuarit të Thelluar, LSTM dhe GRU, rezultuan me gabime më të larta dhe vlera më të ulëta të R^2 , duke treguar performancë më të dobët për dataset-in dhe konfigurimin aktual të sistemit. Duke marrë në konsideratë kompleksitetin e modeleve, saktësinë e parashikimit dhe stabilitetin e rezultateve, modelet XGBoost dhe Regresioni Linear rezultojnë alternativat më të përshtatshme për integrim në sistemin DT dhe modulën e vendimmarrjes automatike për menaxhimin e lagështirës në ambient të brendshëm.

Tabela 6.5 Vlerësimi i modeleve për pragun interval 58-72%

Model	MAE	RMSE	R^2
Regresioni Linear	0.820897066	1.336579378	0.896441
XGBoost	0.777922098	1.366820694	0.891701
Pylli i rastësishëm	1.074898754	1.674496418	0.837457
LSTM	1.906605535	2.280985281	0.698572
GRU	2.311902819	2.631039985	0.598954

Në kuadër të sistemit të vendimmarrjes automatike, kontrolli i mbi-përshtatjes është i rëndësishëm për të siguruar që modelet parashikuese të funksionojnë në mënyrë të qëndrueshme edhe mbi të dhëna të reja të sensorëve. Për këtë arsye, performanca e secilit model u analizua në train dhe test për të verifikuar aftësinë e përgjithësimit të modelit. Rezultatet tregojnë se ndryshimi ndërmjet pragut fiks dhe intervalit operacional nuk ndikon ndjeshëm në performancën regressive të modeleve, pasi metrikat MAE, RMSE dhe R^2 mbeten pothuajse të pandryshuara. Kjo konfirmon se modelet parashikuese ruajnë stabilitetin e tyre në parashikimin e lagështirës për horizontin 30-minutësh.

Figura 6.13 paraqet krahasimin ndërmjet vlerave reale të lagështirës dhe vlerave të parashikuara nga modeli Linear Regression përgjatë intervalit kohor të analizuar. Vërehet se modeli arrin të ndjekë në mënyrë relativisht të mirë trendin e përgjithshëm të të dhënave reale, duke reflektuar si periudhat e rritjes ashtu edhe ato të uljes së lagështirës. Diferencat ndërmjet kurbës reale dhe asaj të parashikuar janë në përgjithësi të vogla, çka tregon se modeli ka aftësi të mirë përshtatjeje ndaj të dhënave. Megjithatë, në disa intervale kohore vërehen devijime të lehta dhe një zbutje e variacioneve të menjëhershme, duke treguar se modeli linear ka kufizime në kapjen e ndryshimeve shumë dinamike ose jolineare të lagështirës.

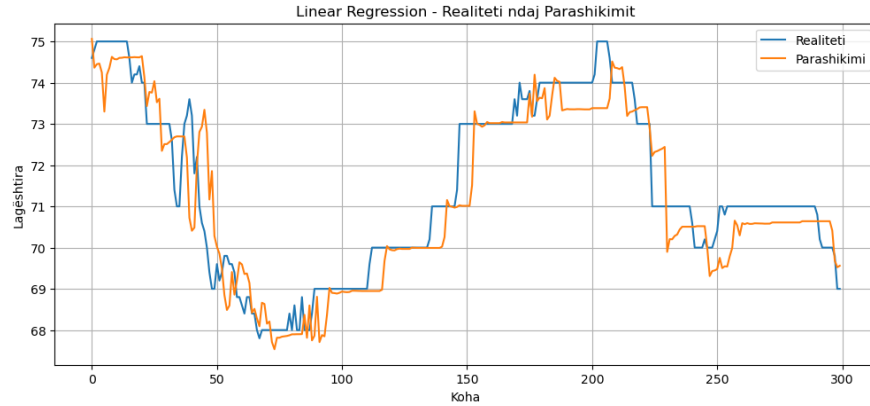


Figura 6.13 Krahاسimi ndërmjet vlerave reale dhe parashikimeve të modelit Linear Regression për lagështirën relative.

Figura 6.14 paraqet performancën e modelit Pyllit të rastësishëm në parashikimin e lagështirës relative krahasuar me vlerat reale. Vërehet se modeli ndjek në mënyrë të kënaqshme trendin kryesor të të dhënave, duke reflektuar shumicën e rritjeve dhe uljeve të lagështirës. Megjithatë, në disa intervale kohore shfaqen devijime më të theksuara dhe luhajtje më të mëdha krahasuar me vlerat reale, çka tregon se modeli është më sensitiv ndaj variacioneve lokale të dataset-it.

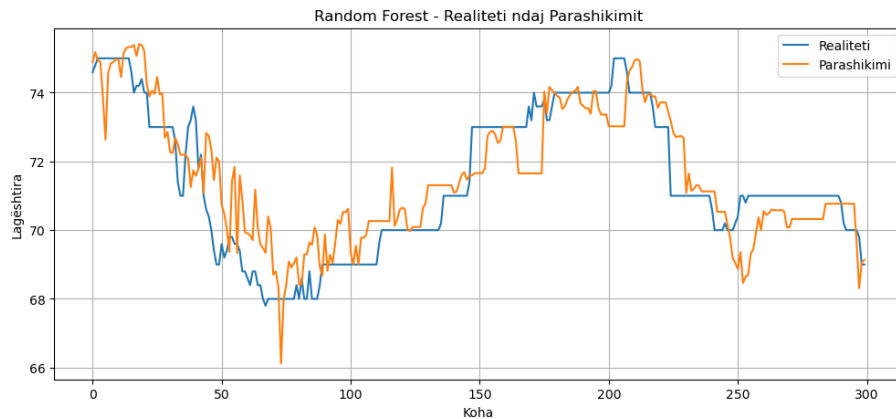


Figura 6.14 Krahاسimi ndërmjet vlerave reale dhe parashikimeve të modelit Pyllit të rastësishëm për lagështirën relative.

Figura 6.15 paraqet krahasimin ndërmjet vlerave reale dhe atyre të parashikuara nga modeli XGBoost. Vërehet se modeli arrin të ndjekë me saktësi trendin e përgjithshëm të lagështirës, duke ruajtur afërsi të mirë me të dhënat reale në shumicën e intervaleve kohore. Diferencat ndërmjet parashikimit dhe realitetit janë relativisht të vogla, çka tregon aftësi të mirë të modelit për të kapur ndryshimet e lagështirës dhe për të realizuar parashikime të qëndrueshme.

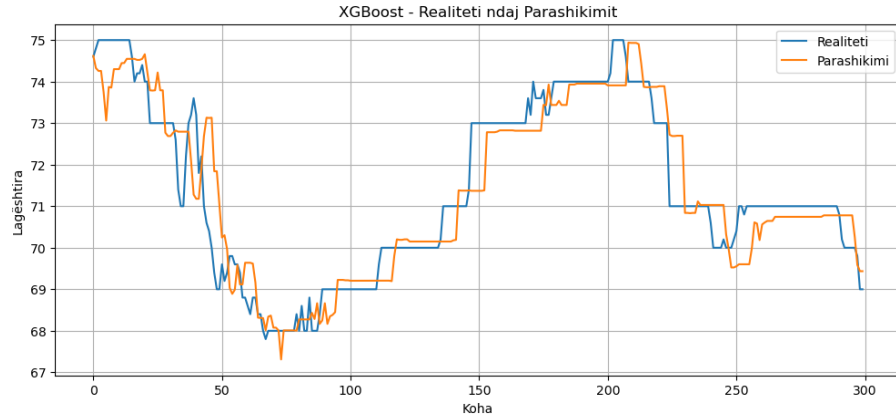


Figura 6.15 Krahasoni ndërmjet vlerave reale dhe parashikimeve të modelit XGBoost për lagështirën relative.

Tabela 6.6 paraqet krahasimin e performancës së modeleve Regresioni Linear, Pylli i Rastësishëm dhe XGBoost në dataset-in e trajnimit dhe testimit duke përdorur metrikat MAE, RMSE dhe R^2 . Analiza e diferencës ndërmjet rezultateve të trainimit dhe testimit shërben për të vlerësuar praninë e mbipërshtatjes (overfitting). Vërehet se modeli Pylli i Rastësishëm arrin performancë shumë të lartë në trainim, por shfaq rënie të konsiderueshme në testim, duke rezultuar në diferencën më të madhe të RMSE (1.056), çka tregon tendencë më të lartë për mbipërshtatje. Në krahasim me të, modeli XGBoost paraqet diferencën më të vogël ndërmjet trainimit dhe testimit (0.425), duke treguar aftësi më të mirë përgjithësimi dhe stabilitet më të lartë ndaj të dhënave të reja. Regresioni Linear paraqet gjithashtu sjellje relativisht të qëndrueshme, por me performancë më të moderuar krahasuar me XGBoost. Si përfundim, rezultatet sugjerojnë se modeli XGBoost ofron balancën më të mirë ndërmjet saktësisë së parashikimit dhe kontrollit të mbipërshtatjes për këtë dataset.

Tabela 6.6 Krahasoni i performancës së modeleve të ML në dataset-in e trajnimit dhe testimit për kontrollin e mbipërshtatjes (overfitting).

	MAE		RMSE		R2		Diferenca RMSE (për mbipërshtatje)
	Trajnim	Testim	Trajnim	Testim	Trajnim	Testim	
Regresioni Linear	1.159	0.821	1.938	1.337	0.980	0.896	0.601
Pylli rastesishem	0.345	1.075	0.618	1.674	0.998	0.837	1.056
XGBoost	1.093	0.778	1.792	1.367	0.983	0.892	0.425

6.6. Autonomia e modulit të vendimarrjes

Moduli i vendimarrjes përfaqëson komponentin logjik të sistemit të automatizuar, i cili analizon periodikisht të dhënat e lagështirës së mjedisit dhe përcakton veprimin e nevojshëm për kontrollin e thithësit të lagështirës. Moduli përditëson çdo 5 minuta një skedar Excel duke shtuar një rresht në fund nga matësi i lagështirës dhe fshin një rresht në

fillim. Në çdo hap kohor, sistemi lexon vlerën aktuale të lagështirës, e krahason atë me pragjet e përcaktuara të kontrollit dhe gjeneron një vendim virtual, si ndezje, fikje ose ruajtje të gjendjes ekzistuese të pajisjes.

Në këtë studim, pragu fiks i lagështirës u përdor fillimisht si skenar krahasues për vlerësimin e modeleve përmes matricës së konfuzionit dhe metrikave Saktësi, Precizioni, Ndjeshmëria dhe Madhësisë F1. Megjithatë, për logjikën finale të vendimmarrjes në sistemin DT u zgjodh qasja me interval kontrolli dhe histerezi, konkretisht 58–72%. Kjo qasje është më e përshtatshme për një proces dinamik si lagështira, pasi shmang ndezjet dhe fikjet e shpeshta të pajisjes dhe rrit stabilitetin operacional të sistemit.

$$u(t) = \begin{cases} 1, & H(t) > 72 \\ 0, & H(t) < 58 \\ u(t-1), & 58 \leq H(t) \leq 72 \end{cases} \quad (6.5)$$

ku:

- $u(t)$ = gjendja e pajisjes (ON/OFF)
- $H(t)$ = lagështira

Në këtë kuadër, modeli XGBoost u konsiderua më i përshtatshëm për integrim në modulin parashikues, pasi ofroi performancë të lartë dhe diferencë më të vogël ndërmjet trainimit dhe testimit, duke treguar aftësi më të mirë përgjithësimi.

Për të vlerësuar performancën e modeleve të inteligjencës artificiale në kontekstin e vendimmarrjes automatike, u përdor matrica e konfuzionit dhe metrika të derivuara prej saj, si Recall, F1-score, True Positive (TP), True Negative (TN), False Positive (FP) dhe False Negative (FN). Këto metrika lejojnë analizimin jo vetëm të saktësisë së përgjithshme të modelit, por edhe ndikimin operacional të vendimeve të marra nga sistemi DT për kontrollin e lagështirës.

Në këtë kontekst, Tabela 6.7 paraqet rastet True Positive që përfaqësojnë aktivizime korrekte të pajisjes kur lagështira tejkalon prapun kritik, ndërsa False Positive tregojnë aktivizime të panevojshme që mund të sjellin konsum shtesë energjie. Nga ana tjetër, False Negative konsiderohen rastet më kritike, pasi sistemi nuk aktivizon pajisjen edhe pse kushtet reale kërkojnë ndërhyrje. Për këtë arsye, analiza e matricës së konfuzionit dhe e metrikave përkatëse është thelbësore për të vlerësuar besueshmërinë e modeleve ML dhe DL në mbështetje të vendimmarrjes automatike.

Tabela 6.7 Përshkrimi i rasteve të vendimmarrjes

Raste e vendimmarrjes	Parashikimi (Duhet aktivizuar)	Veprimi (Është aktivizuar)	Interpretim mbi sistemin
True-Positive (TP)	E vërtetë	E vërtetë	Aktivizim korrekt
True-Negative (TN)	E vërtetë	E gabuar	Mos veprim korrekt
False-Positive (FP)	E gabuar	E vërtetë	Aktivizim i panevojshëm (impakt ekonomik)
False-Negative (FN)	E gabuar	E gabuar	Mos veprim kur është e nevojshme (kritike)

Rëndësi të madhe ka krijimi i kolonës për gjendjen e pajisjes të cilat ruhet gjendja e pajisjes ON/OFF (ndezur ose fikur), krijimi i kolonës së vendimit ku ruhet vendimi i modulit me 4 mundësitë Turn_on, Turn_off, Keep_on, Keep_off (Tabela 6.8)

Tabela 6.8 Lidhja logjike e gjendjes aktuale me atë vendimarrjes

Gjendja aktuale e pajisjes	Lagështira aktuale	Kushti	Vendimi	Gjendja e re
OFF	>72	Kalon pragun e ndezjes	turn_on	ON
OFF	<=72	Nuk kalon pragun e ndezjes	keep_off	OFF
ON	≤ 58	Kalon pragun e fikjes	turn_off	OFF
ON	≥ 58	Nuk kalon pragun e fikjes	keep_on	ON

Një tjetër kolonë e shtuar është kolona e arsyes ku do ruhet shpjegimi logjik i vendimit dhe vlen sidomos për analizë shkencore. Çdo gjë fillon ku gjendja fillestare e pajisjes është e fikur dhe duke përfshirë në cikël të gjithë datasetin plotëson kolonat e shtuara duke marrë për bazë vlerën aktuale të lagështirës. Rastet paraqiten në tabelën e mëposhtme dhe janë: Rasti kur pajisja është e fikur (OFF) dhe lagështira është rritur mbi pragun 72 vendimi i marrë është turn_on. Rasti i dytë është kur lagështira nuk e kalon pragun dhe pajisje mbetet e fikur. Rasti 3, pajisja është e ndezur (statusi ON), lagështira bie nën pragun e fikjes 58 dhe pajisja fiket (kalon në statusin OFF) dhe rasti 4 ku pajisja është e ndezur dhe lagështira është ende mbi pragun e fikjes atëherë vazhdon gjendja ON për pajisjen. Figura 6.16 paraqet strukturën e tabelës së vendimeve të ruajtur në SQLite.

Temperatura	Lageshtira	Trysnia	collected_at	Device_State	Decision	Reason
Filter	Filter	Filter	Filter	Filter	Filter	Filter
16.6	78.0	100234.0	2026-05-10 18:37:28.186580	ON	keep_on	Lagështira 78.00% është mbi pragun e...
16.6	78.0	100236.4	2026-05-10 18:37:28.186580	ON	keep_on	Lagështira 78.00% është mbi pragun e...
16.6	78.0	100240.2	2026-05-10 18:37:28.186580	ON	keep_on	Lagështira 78.00% është mbi pragun e...
16.6	78.0	100240.4	2026-05-10 18:37:28.186580	ON	keep on	Lagështira 78.00% është mbi pragun e...

Figura 6.16 SQLite tabela e vendimeve

Vlera 69 % u konsiderua fillimisht si vlera optimale sepse ndër 98721 matje të regjistruara 48287 janë nën pragun 69 dhe 50434 janë matje me vlerë më të mëdha ose të barabarta me pragun. Performanca e modeleve u vlerësua duke përdorur metrikat si Precizioni (Precision), Ndjeshmëria (Recall), Saktësia e përgjithshme (Accuracy) dhe Masa F1 (F1-score).

Figura 6.17 paraqet krahasimin e metrikave Ndjeshmërisë (Recall) dhe Madhësia F1 për modelet e Regresionit Linear, XGBoost, Pyllit të rastësishëm, LSTM dhe GRU në dy skenarë të vendimarrjes: prag fikse dhe interval me histerezine. Vërehet se modelet tradicionale të ML, veçanërisht XGBoost dhe Regresioni Linear, ruajnë performancë shumë të lartë në të dyja metrikat edhe në rastin e intervalit të kontrollit. Nga ana tjetër, modelet LSTM dhe GRU shfaqin rënie më të ndjeshme të Ndjeshmëri dhe Masën F1 në rastin e intervalit, duke treguar ndjeshmëri më të lartë ndaj ndryshimit të logjikës së vendimarrjes. XGBoost paraqitet si modeli më i qëndrueshëm, duke ruajtur balancë të mirë në të dy skenarët e testuar.

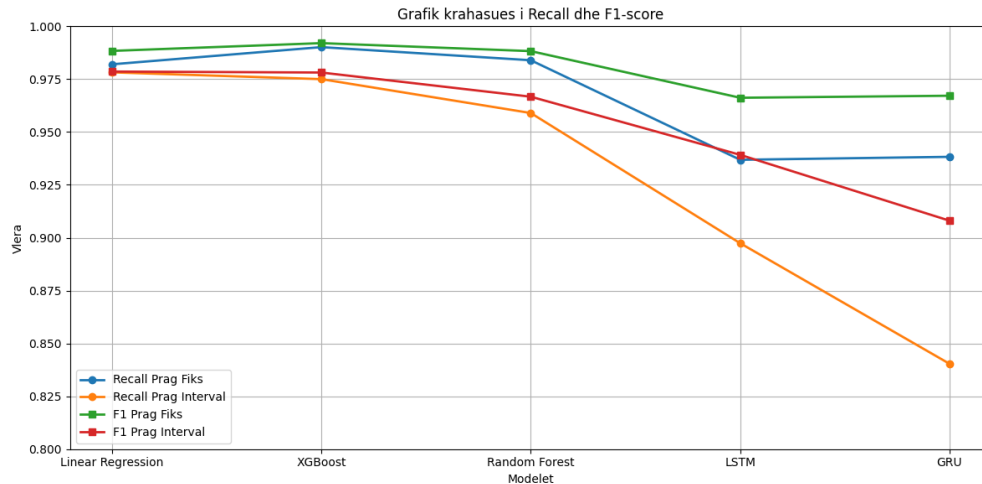


Figura 6.17 Krahasimi i Recall dhe F1-score për modelet parashikuese në rastin e pragut fiks dhe intervalit të vendimmarrjes.

Figura 6.18 paraqet krahasimin e metrikave Saktësisë dhe Precizionit për modelet e testuara në skenarin me prag fiks dhe në atë me interval kontrolli. Rezultatet tregojnë se modelet ML tradicionale ruajnë performancë më të qëndrueshme krahasuar me modelet e të mësuarit të thelluar. XGBoost arrin vlera shumë të larta të Saktësi dhe Precizion si në pragun fiks ashtu edhe në interval, duke treguar aftësi të mirë përgjithësimi dhe stabilitet në vendimmarrje. Ndërkohë, modelet LSTM dhe GRU paraqesin rënie më të konsiderueshme të Saktësinë e përgjithshme në rastin e intervalit, megjithëse ruajnë Precizion relativisht të lartë. Kjo tregon se modelet DL arrijnë të identifikojnë me saktësi rastet pozitive, por performojnë më dobët në stabilitetin e përgjithshëm të klasifikimit.

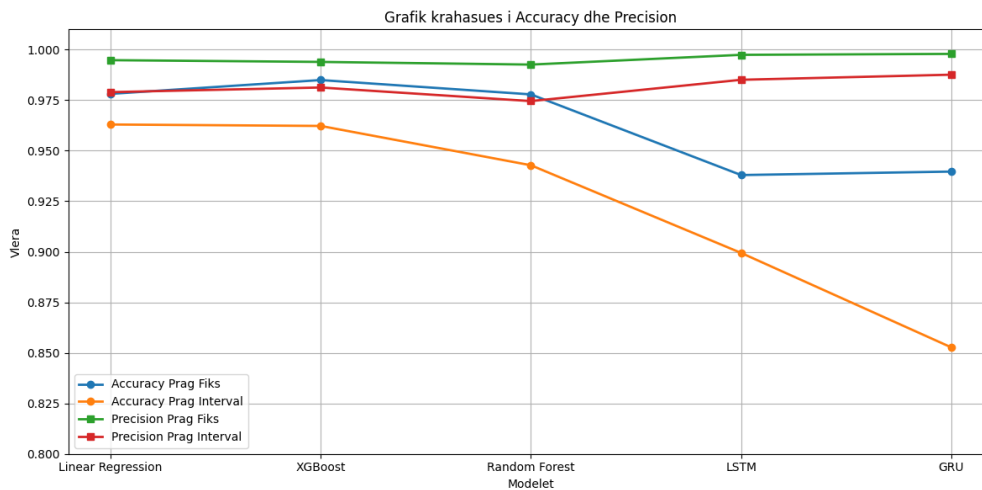


Figura 6.18 Krahasimi i Saktësisë dhe Precizionit për modelet parashikuese në rastin e pragut fiks dhe intervalit të vendimmarrjes.

Duke analizuar në mënyrë të kombinuar metrikat e mësipërme, si dhe stabilitetin ndërmjet pragut fiks dhe intervalit të kontrollit, modeli XGBoost rezulton modeli më i përshtatshëm për integrim në sistemin DT si ndihmesë për realizimin e vendimmarrjes automatike. Ky model ruan performancë shumë të lartë në të gjitha metrikat dhe paraqet ndryshime

minimale ndërmjet skenarëve të testuar, duke treguar aftësi të mirë përgjithësimi dhe mungesë të theksuar të mbi-përshtatjes. Për këtë arsye, XGBoost konsiderohet zgjidhja optimale për parashikimin e lagështirës dhe mbështetjen e vendimmarrjes automatike në kohë reale.

Përveç treguesve kryesor të performancës, si Saktësia, Precizioni, Ndjeshmëria dhe Masa F1, në analizën e secilit model u morën në konsideratë edhe vlerat e matricës së konfuzionit, përfshirë True Positive (TP), True Negative (TN), False Positive (FP) dhe False Negative (FN). Studimi i këtyre elementeve është thelbësor në sistemet e vendimmarrjes automatike, pasi ato reflektojnë mënyrën konkrete se si sistemi Digital Twin reagon dhe merr vendime në skenarë të ndryshëm operacionalë.

Tabela 6.9 Krahasimi i rasteve TP, TN, FP dhe FN për modelet parashikuese

		TP	TN	FP	FN		TP	TN	FP	FN
Regresioni Linear	Prag interval	16716	2294	359	372	Prag fiks	18351	958	96	336
XGBoost		16661	2335	318	427		18503	941	113	184
Pylli rastësishëm		16387	2225	428	701		18387	917	137	300
LSTM		15324	2421	232	1752		17496	1009	45	1179
GRU		14350	2473	180	2726		17522	1017	37	1153

Tabela 6.9 paraqet krahasimin e komponentëve të matricës së konfuzionit për modelet Regresionit Linear, XGBoost, Pyllit të rastësishëm, LSTM dhe GRU në dy skenarë të vendimmarrjes: prag si interval kontrolli dhe prag fiks. Analiza e rasteve TP, TN, FP dhe FN është veçanërisht e rëndësishme në kontekstin e sistemit DT, pasi këto vlera reflektojnë drejtpërdrejt sjelljen operationale të sistemit të kontrollit të lagështirës. Vërehet se përdorimi i pragut fiks rrit numrin e rasteve TP dhe redukton ndjeshëm rastet FN për shumicën e modeleve, duke përmirësuar aftësinë e sistemit për të identifikuar situatat kritike që kërkojnë aktivizimin e pajisjes. Modeli XGBoost paraqet performancën më të balancuar, duke arritur numër të lartë TP dhe numrin më të ulët të rasteve FN në skenarin me prag fiks, çka tregon aftësi më të mirë për të shmangur mosaktivizimet kritike të pajisjes. Nga ana tjetër, modelet LSTM dhe GRU paraqesin numër më të lartë të rasteve FN, veçanërisht në rastin e intervalit të kontrollit, duke treguar ndjeshmëri më të lartë ndaj logjikës së histerezisë. Në përgjithësi, rezultatet tregojnë se XGBoost ofron kompromisin më të mirë ndërmjet identifikimit korrekt të rasteve kritike dhe reduktimit të vendimeve të gabuara në sistemin e vendimmarrjes automatike.

Për realizimin e një sistemi eksperimental të vendimmarrjes automatike në menaxhimin e cilësisë së ajrit në ambient të brendshëm, u implementua një konfigurim praktik që mundëson kontrollin automatik të një pajisjeje për reduktimin e lagështirës përmes një prize inteligjente të bazuar në platformën Tuya (Figura 6.19).



Figura 6.19 Thithësi lagështirës dhe priza inteligjente

Objektivi kryesor i këtij konfigurimi ishte ndërtimi i një infrastrukture funksionale që integron të dhënat e sensorëve të mjedisit me një modul inteligjent të vendimmarrjes dhe me komponentët fizikë të sistemit, duke krijuar kështu bazën funksionale të një sistemi Digital Twin (DT).

Në këtë arkitekturë, pajisja kryesore fizike është një thithës lagështirës, i cili nuk kishte funksionalitete të integruara WiFi ose API për kontroll programatik. Për këtë arsye, u përdor një prizë inteligjente kompatibël me ekosistemin Tuya, i cili vepron si ndërmjetës midis sistemit inteligjent dhe pajisjes fizike. Smart plug-u u lidh me rrjetin lokal WiFi dhe u konfigurua për të kontrolluar furnizimin me energji elektrike të thithësit të lagështirës. Kjo qasje lejon ndezjen dhe fikjen automatike të pajisjes përmes komandave të dërguara nga gjuha Python. Procesi i konfigurimit filloi me krijimin e një projekti në platformën cloud për zhvillues, Tuya (figura 6.20).

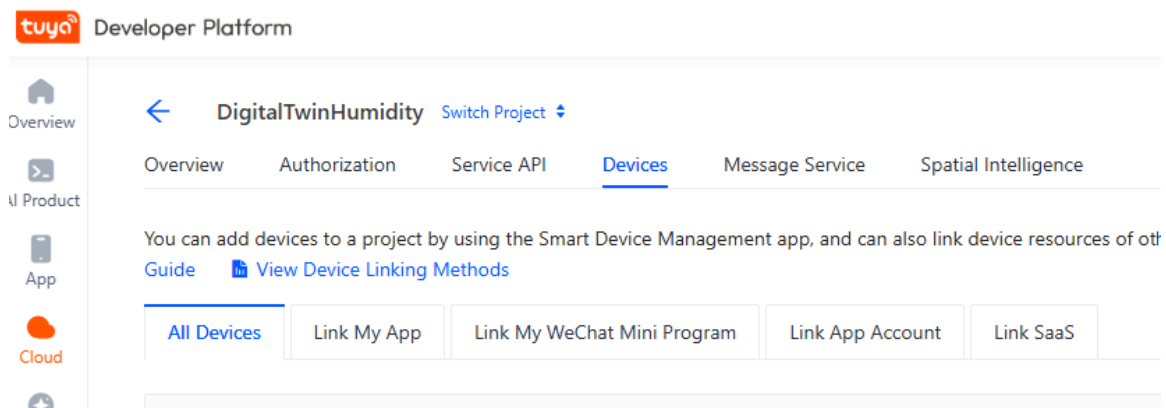


Figura 6.20 Platforma e zhvilluesve Tuya

Në këtë platformë u krijua projekti “DigitalTwinHumidity”, i cili u përdor për menaxhimin e pajisjeve IoT dhe për marrjen e kredencialeve të nevojshme të komunikimit. Gjatë konfigurimit u përdor rajoni “Central Europe Data Center”, i cili përcakton zonën cloud ku regjistrohen pajisjet dhe API-të.

Fillimisht u tentua përdorimi i aplikacionit “Tuya Spatial”, por u konstatua se ky aplikacion nuk ishte i përshtatshëm për menaxhimin standard të smart plug-eve dhe integrimin me TinyTuya. Për këtë arsye u përdor aplikacioni “Smart Life”, i cili është aplikacioni standard i ekosistemit Tuya për pajisje smart home. Smart plug-u u regjistrua dhe u konfigurua në Smart Life duke përdorur rrjetin WiFi 2.4 GHz, i cili zakonisht kërkohet nga shumica e

pajisjeve IoT të kësaj kategorie.

Pas konfigurimit të smart plug-ut në aplikacionin Smart Life, llogaria e përdoruesit u lidh me projektin cloud në Tuya Developer Platform përmes funksionit “Link Tuya App Account”. Ky hap ishte kritik pasi mundësoi sinkronizimin e pajisjeve midis aplikacionit mobil dhe platformës cloud. Pas lidhjes së suksesshme, pajisja smart plug u shfaq në seksionin “Devices” të projektit cloud, duke bërë të mundur aksesin programatik ndaj saj. Më pas u përdor biblioteka TinyTuya në Python për të realizuar komunikimin lokal me pajisjen. TinyTuya është një bibliotekë open-source që lejon kontrollin lokal të pajisjeve Tuya pa pasur nevojë për komunikim të vazhdueshëm me cloud-in. Kjo qasje redukton vonesat në komunikim, shmang varësinë nga interneti dhe eliminon kostot e mundshme të përdorimit intensiv të API-ve cloud.

Fillimisht u instalua biblioteka tinytuya në mjedisin Windows dhe më pas u aktivizua wizard për identifikimin automatikisht të pajisjes në rrjetin lokal dhe për të gjeneruar parametrat e nevojshëm të konfigurimit. Gjatë këtij procesi u përdorën kredencialet e marra nga Tuya Cloud, përfshirë:

- Access ID,
- Access Secret,
- region “eu”.

Wizard-i realizoi skanimin e rrjetit lokal dhe identifikoi prizën inteligjente me adresë lokale IP. Gjithashtu u gjeneruan parametrat: Device ID, Local Key, IP Address, versioni i protokollit të komunikimit.

Këta parametra u ruajtën automatikisht në skedarin devices.json, i cili përdoret më pas nga kodi Python për komunikimin me pajisjen.

Pas konfigurimit, komunikimi lokal me prizën inteligjente u realizua përmes klasës OutletDevice të TinyTuya. Në këtë fazë u testuan disa versione të protokollit të komunikimit (3.3, 3.4 dhe 3.5), pasi pajisje të ndryshme Tuya përdorin versione të ndryshme të protokollit. Gjatë testimit u konstatua se komanda kryesore për ndezje/fikje ishte e lidhur me DPS (Data Point Service) numër 1, ku:

- përgjigjia True përfaqësonte ndezjen,
- përgjigjia False përfaqësonte fikjen e pajisjes.

Kjo infrastrukturë krijon bazën funksionale të Digital Twin-it, ku modeli virtual merr të dhëna nga sensorët e lagështirës, realizon analizë dhe vendimmarrje automatike dhe më pas dërgon komanda reale drejt pajisjes fizike. Në praktikë, sistemi mund të funksionojë sipas një politike të thjeshtë kontrolli, ku dehumidifier-i ndizet automatikisht kur lagështira tejkalon një prag të caktuar, për shembull 60%, dhe fiket kur lagështira rikthehet në nivele të pranueshme.

Përfundime

Në këtë disertacion u trajtua në mënyrë të integruara roli i teknikave të Inteligjencës Artificiale në ndihmesë të realizimit të vendimarrjes automatike, duke kombinuar analizën teorike me implementime praktike në tre raste studimore të ndryshme: chatbot-et inteligjente me XAI në hoteleri, analiza e sentimentit nga e folura në gjuhën shqipe dhe sistemi i bazuar në DT (Digital Twin) për menaxhimin e lagështirës në ambiente të brendshme. Qasja e ndjekur tregoi se vendimarrja automatike nuk duhet parë vetëm si një proces teknik për të realizuar optimizim, por si një mekanizëm i plotë që përfshin etapa si mbledhja e të dhënave, analiza, parashikimi i rezultateve si dhe interpretimin dhe ekzekutimin e vendimeve në kohë reale.

Nga analiza teorike u evidentua se evolucioni i Inteligjencës Artificiale ka transformuar proceset tradicionale të vendimarrjes nga sisteme të ngurta të bazuara në rregulla fikse drejt sistemeve inteligjente dhe adaptive, të cilat analizojnë, mësojnë nga të dhënat dhe funksionojnë në mjedise dinamike e komplekse. Studimi konfirmoi se modelet moderne të ML dhe DL mund të përmirësojnë ndjeshëm cilësinë, shpejtësinë dhe shkallëzimin e vendimeve, duke transformuar vendimarrjen nga një qasje që reagon ndaj ngjarjeve, në një qasje parashikuese dhe proaktive që mundëson ndërhyrje të hershme dhe më inteligjente. Në të njëjtën kohë, u evidentuan edhe sfidat që lidhen me mungesën e transparencës dhe interpretueshmërisë së modeleve komplekse, veçanërisht të atyre të bazuara në DL (Deep Learning) dhe arkitektura Transformer.

Në rastin e parë studimor, u demonstrua se integrimi i metodave të XAI në chatbot-et inteligjente mund të rrisë transparencën dhe besueshmërinë e sistemeve vendimarrëse në hoteleri. Përdorimi i teknikave interpretuese si SHAP dhe analiza e funksionimit të modelit tregoi se vendimet e gjeneruara nga sistemi mund të bëhen më të kuptueshme për përdoruesit dhe operatorët njerëzorë. Kjo konfirmoi hipotezën se sistemet që ndërthurin performancën e lartë me aftësinë për të ofruar shpjegime janë më të besueshme dhe më lehtësisht të pranueshme në skenarë realë, veçanërisht në mjedise ku besimi i përdoruesit luan rol kyç.

Në rastin e dytë studimor, analiza e sentimentit nga e folura në gjuhën shqipe tregoi se modelet e DL mund të funksionojnë me sukses edhe në gjuhë me burime të kufizuara. Procesi i normalizimit të audios, ekstraktimit të veçorive dhe aplikimit të modeleve BiLSTM dhe HuBERT tregoi se klasifikimi i emocioneve mund të realizohet me një nivel të lartë performance, edhe pse vazhdon të ekzistojë një farë mbivendosjeje ndërmjet klasave emocionale. Rezultatet konfirmuan rëndësinë e kombinimit të teknikave tradicionale të veçorive akustike me modelet moderne Transformer për përmirësimin e klasifikimit emocional dhe për ndërtimin e sistemeve automatike të perceptimit njerëzor.

Rasti i tretë studimor, i fokusuar në DT për menaxhimin e lagështirës relative, demonstroi se kombinimi i sensorëve IoT, modeleve të parashikimit dhe mekanizmave të vendimarrjes krijon mundësinë për zhvillimin e sistemeve autonome, të afta për të ndërvepruar dhe reaguar në kohë reale ndaj ndryshimeve të mjedisit. Analiza e strukturës së të dhënave konfirmoi se dataset-i kishte cilësi të mirë për modelim parashikues, ndërsa

krahasimi i modeleve statistikore, ML dhe DL evidentoi se modelet jo-lineare kishin performancë më të mirë në identifikimin e rasteve kritike të lagështirës. Gjithashtu, u demonstrua se vendimmarrja automatike mund të realizohet përmes kombinimit të pragjeve logjike me parashikimet e modeleve, duke krijuar një mekanizëm funksional për kontrollin automatik të pajisjeve fizike në një mjedis DT.

Në përmbledhje, ky disertacion konfirmoi se teknikat e inteligjencës artificiale mund të përdoren në mënyrë efektive për realizimin e vendimmarrjes automatike në kontekste heterogjene dhe me tipologji të ndryshme të dhënash. Studimi tregoi se inteligjenca artificiale nuk kufizohet vetëm në analizën e të dhënave, por mund të shndërrohet në një mekanizëm të plotë vendimmarrës që kombinon perceptimin, analizën, parashikimin dhe veprimin automatik. Njëkohësisht, rezultatet theksuan se transparenca, interpretueshmëria dhe kontrolli njerëzor mbeten elemente të domosdoshme për ndërtimin e sistemeve të besueshme dhe të pranueshme në praktikë. Për këtë arsye, kontributi kryesor i këtij punimi qëndron në demonstrimin e faktit se performanca e lartë dhe shpjegueshmëria nuk janë koncepte kundërshtuese, por mund të integrohen në mënyrë të balancuar për ndërtimin e sistemeve inteligjente të qëndrueshme, adaptive dhe të përgjegjshme.

Punime për të ardhmen

Në punime të ardhshme, sistemi i chatbot-it inteligjent me XAI mund të zgjerohet drejt ndërtimit të një platforme multimodale dhe adaptive për vendimmarrje në hoteleri. Një drejtim i rëndësishëm është integrimi i modeleve të avancuara LLM (Large Language Models) me mekanizma të memories kontekstuale dhe profileve të përdoruesve, në mënyrë që sistemi të ofrojë rekomandime dhe përgjigje më të personalizuar. Gjithashtu, mund të implementohen mekanizma të vazhdueshëm të të mësuarit, ku sistemi përshtatet dinamikisht bazuar në ndërveprimet me klientin. Ndërtimi i një modeli hibrid ku kombinohen metoda të ndryshme shpjeguese si SHAP, LIME dhe Vizualizimi Vëmendjes për të ofruar shpjegime më intuitive dhe më të kuptueshme për përdoruesit dhe operatorët e hotelit. Në këtë mënyrë, sistemi nuk do të japë vetëm rekomandime automatike, por edhe arsyetimin logjik pas çdo vendimi.

Një drejtim i rëndësishëm për punime të ardhshme është zgjerimi i analizës së sentimentit nga e folura përmes integrit të një qasjeje multimodale që kombinon sinjalin audio me analizën tekstuale në gjuhën shqipe. Në këtë kontekst, audio do të përpunohet fillimisht përmes teknikave të njohjes automatike të të folurit (Automatic Speech Recognition – ASR), duke realizuar transkriptimin automatik të të folurës në tekst. Më pas, si karakteristikat akustike ashtu edhe përmbajtja tekstuale do të analizohen paralelisht përmes modeleve të inteligjencës artificiale dhe të mësuarit të thellë. Kjo qasje pritet të përmirësojë saktësinë e njohjes së emocioneve dhe sentimentit në gjuhën shqipe, duke kombinuar informacionin prosodik dhe tonal të të folurit me informacionin semantik dhe kontekstual të tekstit. Gjithashtu, integrimi i analizës multimodale mund të kontribuojë në ndërtimin e sistemeve më të avancuara dhe më të besueshme të vendimmarrjes automatike në aplikime reale të ndërveprimit njeri–makinë.

Në punime të ardhshme, framework-u i propozuar i Binjakut Dixhital mund të zgjerohet përmes integrit të teknikave të RL (Reinforcement Learning), me qëllim realizimin e një sistemi më autonom dhe adaptiv të vendimmarrjes automatike. Ndryshe nga qasja aktuale e bazuar në pragje dhe modele parashikuese statike, RL do t'i mundësonte sistemit të mësojë politika optimale kontrolli përmes ndërveprimit të vazhdueshëm me ambientin. Në këtë kontekst, agjenti inteligjent do të analizonte gjendjen aktuale të mjedisit të brendshëm, si lagështira, temperatura dhe parametrat e tjerë IAQ, dhe do të merrte vendime automatike për aktivizimin ose çaktivizimin e pajisjeve, duke optimizuar njëkohësisht komfortin e përdoruesit dhe konsumin energjetik. Përmes mekanizmit të shpërblimit , sistemi mund të mësojë gradualisht strategji më efikase dhe më adaptive krahasuar me rregullat fikse tradicionale.

Gjithashtu, kombinimi i DT me RL mund të mundësojë simulimin e skenarëve të ndryshëm mjedisorë dhe testimin e politikave vendimmarrëse në mënyrë virtuale përpara aplikimit në sistemin real fizik. Kjo qasje do të rrisë autonominë, fleksibilitetin dhe inteligjencën adaptive të sistemit të propozuar.

Aneks -Kodi implementuar

Analiza e të dhënave audio për dy karakteristikat kryesore akustike, rrafshësia spektrale (Spectral flatness) dhe shkalla e kalimit nëpër zero (zero crossing rate) për të dalluar sinjale me strukturë harmonike nga zhurma dhe frekuencën e kalimit të sinjalit nëpër zero. Vendoset pragjet empirike për të identifikuar regjistrimet me zhurmë.

```
import os
import librosa
import numpy as np
import matplotlib.pyplot as plt

def get_features(file_path):
    y, sr = librosa.load(file_path, sr=None, mono=True)
    flatness = float(librosa.feature.spectral_flatness(y=y).mean())
    zcr = float(librosa.feature.zero_crossing_rate(y).mean())
    return flatness, zcr

folder_path = r"/content/drive/MyDrive/zeri/neutral"

X_flat, Y_zcr, names = [], [], []
for fname in os.listdir(folder_path):
    if fname.lower().endswith(".wav"):
        fpath = os.path.join(folder_path, fname)
        try:
            f, z = get_features(fpath)
            X_flat.append(f); Y_zcr.append(z); names.append(fname)
        except Exception as e:
            print(f"[ERROR] {fname}: {e}")

X_flat = np.array(X_flat)
Y_zcr = np.array(Y_zcr)

flat_thresh = 0.40
zcr_thresh = 0.05
is_noise = (X_flat > flat_thresh) & (Y_zcr > zcr_thresh)

plt.figure(figsize=(8,6))

plt.scatter(X_flat[~is_noise], Y_zcr[~is_noise], alpha=0.7, edgecolors='k', label="jo-zhurmë")

plt.scatter(X_flat[is_noise], Y_zcr[is_noise], marker='x', s=100, linewidths=2, label="zhurmë")

plt.axvline(flat_thresh, linestyle="--", linewidth=1)
plt.axhline(zcr_thresh, linestyle="--", linewidth=1)

plt.xlabel("Rrafshësia spektrale (mean)")
plt.ylabel("Shkalla e kalimit nëpër zero (mean)")
plt.title("Rrafshësia ndaj Shkallës ")
plt.grid(True)
plt.legend()
plt.tight_layout()
plt.show()
```

Procesi i përgatitjes së dataset-it audio që të jetë i përshtatshëm për modelet ML/DL. Konvertohen në mono, normalizohet amplituda, reduktohet zhurma dhe hiqet heshtja, segmentimi në pjesë më të vogla, lidhja e çdo audio me etiketën dhe krijohet skedari metadata për trajnim dhe analizë

```
import argparse, os, sys
import numpy as np
import pandas as pd
import librosa, soundfile as sf
import noisereduce as nr

# ----- Funksionet ndihmëse -----
def peak_normalize(x, eps=1e-9):
    peak = np.max(np.abs(x)) + eps
    return x / peak

def to_zero_one(x):
    return (x + 1.0) / 2.0

def remove_silence(y, sr, top_db=30, min_interval_ms=100, pad_ms=20):
    intervals = librosa.effects.split(y, top_db=top_db)
    if len(intervals) == 0:
        return y
    min_samples = int((min_interval_ms/1000.0) * sr)
    pad = int((pad_ms/1000.0) * sr)
    chunks = []
    for (s, e) in intervals:
        if e - s >= min_samples:
            s2 = max(0, s - pad)
            e2 = min(len(y), e + pad)
            chunks.append(y[s2:e2])
    if len(chunks) == 0:
        return y
    return np.concatenate(chunks)

def apply_denoise(y, sr, noise_seconds=0.5):
    if len(y) < int(noise_seconds * sr):
        return nr.reduce_noise(y=y, sr=sr)
    noise_profile = y[: int(noise_seconds * sr)]
    return nr.reduce_noise(y=y, y_noise=noise_profile, sr=sr)

def resample_audio(y, orig_sr, target_sr):
    if orig_sr == target_sr:
        return y, orig_sr
    y_rs = librosa.resample(y=y, orig_sr=orig_sr, target_sr=target_sr, res_type="soxr_hq")
    return y_rs, target_sr

def segment_audio_with_times(y, sr, chunk_seconds=2.0, overlap_seconds=0.0):
    """Kthen (segments, start_times_sec, end_times_sec)."""
    win = int(chunk_seconds * sr)
    hop = int(max(0.001, (chunk_seconds - overlap_seconds)) * sr)

    if len(y) <= win:
        seg = y
        # pad nëse do gjatësi fikse
```

```

    if len(seg) < win:
        seg = np.pad(seg, (0, win-len(seg)), mode='constant')
    return [seg], [0.0], [min(len(y)/sr, chunk_seconds)]

segs, starts_sec, ends_sec = [], [], []
start = 0
while start + win <= len(y):
    segs.append(y[start:start+win])
    starts_sec.append(start / sr)
    ends_sec.append((start + win) / sr)
    start += hop

# Opsionale: ruaj bishtin në fund nëse >= 0.5s
tail = y[start:]
if len(tail) >= int(0.5 * sr):
    pad = win - len(tail)
    seg = np.pad(tail, (0, pad), mode='constant')
    segs.append(seg)
    starts_sec.append(start / sr)
    ends_sec.append((start + win) / sr)

return segs, starts_sec, ends_sec
# ----- fundi -----
def process_file(in_path, args, out_dir, label_map, meta_rows):
    base = "/content/drive/MyDrive/zeri/positive" # os.path.splitext(os.path.basename(in_path))[0]
    y, sr0 = librosa.load(in_path, sr=None, mono=True)
    y = peak_normalize(y)
    y = remove_silence(y, sr0, top_db=float(args.top_db))
    if args.denoise:
        y = apply_denoise(y, sr0)
    target_sr = args.target_sr if args.target_sr is not None else (16000 if args.mode in ["hubert", "mfcc"] else
16000)
    y, sr = resample_audio(y, sr0, target_sr)
    y = peak_normalize(y)
    y_out = to_zero_one(y) if args.normalize_01 else y

# Segment me kohë
segments, starts_sec, ends_sec = segment_audio_with_times(
    y_out, sr, chunk_seconds=args.chunk_seconds, overlap_seconds=args.overlap_seconds
)

# Merr etiketën nga label_map (bazuar në emrin e file-it)
src_name = os.path.basename(in_path)
label = label_map.get(src_name) # mund të jetë None nëse s'është në Excel

# Save segments + mbledh metadata rresht për rresht
for i, (seg, t0, t1) in enumerate(zip(segments, starts_sec, ends_sec), start=1):
    out_path = os.path.join(out_dir, f"{base}_seg{i:04d}.wav")
    to_save = (seg * 2 - 1.0) if args.normalize_01 else seg
    sf.write(out_path, to_save, sr, subtype="PCM_16")

meta_rows.append({
    "segment_path": out_path,
    "source_file": src_name,
    "label": label,
    "segment_index": i,

```

```

        "start_sec": round(float(t0), 4),
        "end_sec": round(float(t1), 4),
        "sr": sr,
        "chunk_seconds": args.chunk_seconds,
        "overlap_seconds": args.overlap_seconds
    })

# Ruaj edhe full_clean
full_clean_path = os.path.join(out_dir, f"{base}_full_clean.wav")
full_to_save = (y_out * 2 - 1.0) if args.normalize_01 else y_out
sf.write(full_clean_path, full_to_save, sr, subtype="PCM_16")

# Rresht meta për full clean (opsionale)
meta_rows.append({
    "segment_path": full_clean_path,
    "source_file": src_name,
    "label": label,
    "segment_index": 0, # 0 = file i plotë
    "start_sec": 0.0,
    "end_sec": round(float(len(y_out)/sr), 4),
    "sr": sr,
    "chunk_seconds": None,
    "overlap_seconds": None
})

print(f"U procesua: {in_path} | Label: {label} | Segments: {len(segments)}")

def main():
    folder = "/content/drive/MyDrive/zeri/positive"
    out_dir = "/content/drive/MyDrive/zeri/positive_processed_plus"
    labels_xlsx = "/content/drive/MyDrive/zeri/labels.xlsx"

    os.makedirs(out_dir, exist_ok=True)
    # Lexohet Excel-in i etiketiveve
    if not os.path.exists(labels_xlsx):
        print(f"Exceli nuk u gjet: {labels_xlsx}. Do vazhdohet pa etiketa.")
        label_map = {}
    else:
        df = pd.read_excel(labels_xlsx)
        if "path" not in df.columns or "label" not in df.columns:
            raise ValueError("Excel duhet të ketë kolonat: 'path' dhe 'label'")
        df["basename"] = df["path"].apply(lambda p: os.path.basename(str(p)))
        label_map = {bn: lab for bn, lab in zip(df["basename"], df["label"])}

    meta_rows = []
    if not os.path.exists(folder):
        raise FileNotFoundError(f"Nuk ekziston input_folder: {folder}")

    for fname in os.listdir(folder):
        if fname.lower().endswith(".wav"):
            in_path = os.path.join(folder, fname)
            process_file(in_path, None, out_dir, label_map, meta_rows)

    meta_csv_path = os.path.join(out_dir, "segments_metadata.csv")
    pd.DataFrame(meta_rows).to_csv(meta_csv_path, index=False, encoding="utf-8")
    print(f"Metadata u ruajt: {meta_csv_path}")

```

```
print(f"Output audio dir: {out_dir}")
```

Skripti që bë krahasimin mbi të njëjtin dataset trajnimi/testi përdor MFCC dhe HuBERT dhe raporton metrikat (Saktësinë, Madhësinë F1, Matricen e Konfuzionit)

```
MFCC_CSV = "/content/drive/MyDrive/zeri/mfcc/features_mfcc_total.csv"
```

```
HUBERT_CSV = "/content/drive/MyDrive/zeri/hubert_features_total.csv"
```

```
KEY_COL = None
```

```
LABEL_COL = "label"
```

```
MFCC_PREFIX = ("feat")
```

```
HUBERT_PREFIX = ("hubert",)
```

```
TEST_SIZE = 0.20
```

```
VAL_SIZE = 0.20
```

```
RANDOM_SEED = 42
```

```
USE_MLP = True
```

```
import pandas as pd
```

```
import numpy as np
```

```
from sklearn.model_selection import train_test_split
```

```
from sklearn.preprocessing import StandardScaler
```

```
from sklearn.pipeline import Pipeline
```

```
from sklearn.linear_model import LogisticRegression
```

```
from sklearn.neural_network import MLPClassifier
```

```
from sklearn.metrics import accuracy_score, f1_score, classification_report, confusion_matrix
```

```
df_mfcc = pd.read_csv(MFCC_CSV)
```

```
df_hubert = pd.read_csv(HUBERT_CSV)
```

```
def feat_cols(df, prefixes):
```

```
    return [c for c in df.columns if any(c.startswith(p) for p in prefixes)]
```

```
mfcc_cols = feat_cols(df_mfcc, MFCC_PREFIX)
```

```
hubert_cols = feat_cols(df_hubert, HUBERT_PREFIX)
```

```
if KEY_COL is not None:
```

```
    assert KEY_COL in df_mfcc.columns and KEY_COL in df_hubert.columns, \
```

```
        f"Kolona kyç '{KEY_COL}' duhet të jetë në të dy CSV-të."
```

```
    df_mfcc = df_mfcc[[KEY_COL, LABEL_COL] + mfcc_cols]
```

```
    df_hubert = df_hubert[[KEY_COL, LABEL_COL] + hubert_cols]
```

```
    merged = pd.merge(df_mfcc[[KEY_COL, LABEL_COL]],
```

```
                        df_hubert[[KEY_COL, LABEL_COL]], on=KEY_COL, suffixes=("_mfcc", "_hubert"))
```

```
    if not (merged[f"{LABEL_COL}_mfcc"] == merged[f"{LABEL_COL}_hubert"]).all():
```

```
        print("Etiketat nuk përputhen ndërmjet CSV-ve. Do të mbaj etiketen nga MFCC.")
```

```
    df_hubert = pd.merge(df_hubert.drop(columns=[LABEL_COL]),
```

```
                        df_mfcc[[KEY_COL, LABEL_COL]],
```

```
                        on=KEY_COL, how="inner")
```

```
else:
```

```
    pass
```

```
def split_xy(df, feat_columns):
```

```
    X = df[feat_columns].astype(np.float32).values
```

```
    y = df[LABEL_COL].values
```

```
    X_tr, X_tmp, y_tr, y_tmp = train_test_split(
```

```
        X, y, test_size=TEST_SIZE+VAL_SIZE, stratify=y, random_state=RANDOM_SEED
```

```
    )
```

```
    rel_val = VAL_SIZE / (TEST_SIZE + VAL_SIZE)
```

```
    X_va, X_te, y_va, y_te = train_test_split(
```

```
        X_tmp, y_tmp, test_size=1-rel_val, stratify=y_tmp, random_state=RANDOM_SEED
```

```

    )
    return X_tr, y_tr, X_va, y_va, X_te, y_te

Xtr_m, ytr_m, Xva_m, yva_m, Xte_m, yte_m = split_xy(df_mfcc, mfcc_cols)
Xtr_h, ytr_h, Xva_h, yva_h, Xte_h, yte_h = split_xy(df_hubert, hubert_cols)

def evaluate_model(name, Xtr, ytr, Xva, yva, Xte, yte, use_mlp=USE_MLP):
    results = []
    logreg = Pipeline([
        ("scaler", StandardScaler()),
        ("clf", LogisticRegression(max_iter=1000, class_weight="balanced"))
    ])
    logreg.fit(Xtr, ytr)
    for split_name, X, y in [("valid", Xva, yva), ("test", Xte, yte)]:
        yp = logreg.predict(X)
        acc = accuracy_score(y, yp)
        fl = f1_score(y, yp, average="macro")
        results.append([name, "LogReg", split_name, acc, fl])
        if split_name == "test":
            print(f"\n[ {name} | LogReg | TEST]")
            print(classification_report(y, yp, digits=3))
            print("Confusion matrix:\n", confusion_matrix(y, yp))
    if use_mlp:
        mlp = Pipeline([
            ("scaler", StandardScaler()),
            ("clf", MLPClassifier(hidden_layer_sizes=(128,),
                                  activation="relu", alpha=1e-4,
                                  batch_size=64, learning_rate_init=1e-3,
                                  max_iter=200, random_state=RANDOM_SEED))
        ])
        mlp.fit(Xtr, ytr)
        for split_name, X, y in [("valid", Xva, yva), ("test", Xte, yte)]:
            yp = mlp.predict(X)
            acc = accuracy_score(y, yp)
            fl = f1_score(y, yp, average="macro")
            results.append([name, "MLP(128)", split_name, acc, fl])
            if split_name == "test":
                print(f"\n[ {name} | MLP(128) | TEST]")
                print(classification_report(y, yp, digits=3))
                print("Matrica Konfuzionit:\n", confusion_matrix(y, yp))

    return results

res_mfcc = evaluate_model("MFCC", Xtr_m, ytr_m, Xva_m, yva_m, Xte_m, yte_m)
res_hubert = evaluate_model("HuBERT", Xtr_h, ytr_h, Xva_h, yva_h, Xte_h, yte_h)

# ----- Tabela përmbledhëse -----
summary = pd.DataFrame(res_mfcc + res_hubert,
                        columns=["features", "model", "split", "accuracy", "macro_f1"])
print("\n=== SUMMARY ===")
print(summary.pivot_table(index=["features", "model"],
                           columns="split", values=["accuracy", "macro_f1"]).round(3))

```

```
# =====
KRIJIMI MODELIT BiLSTM PËR KLASIFIKIMIN E TË DHËNAVE SEKUENCIALE
# =====
```

```
import torch.nn as nn
class BiLSTMClassifier(nn.Module):
    def __init__(self, feat_dim=5, hidden=128, num_classes=3):
        super().__init__()
        self.lstm = nn.LSTM(feat_dim, hidden, batch_first=True, bidirectional=True)
        self.dropout = nn.Dropout(0.3)
        self.fc = nn.Linear(hidden*2, num_classes)
    def forward(self, x):
        out, _ = self.lstm(x)
        out = out.mean(dim=1)
        out = self.dropout(out)
        return self.fc(out)
model = BiLSTMClassifier(feat_dim=5, hidden=128, num_classes=len(le.classes_))
```

```
# =====
# KRAHASIMI I MODELEVE PËR HORIZONTIN 30 MINUTA
# Regression + Vlerësimi vendimarrjes me interval pragjesh
# =====
```

```
import pandas as pd
import numpy as np
import sqlite3

from sklearn.model_selection import train_test_split
from sklearn.preprocessing import StandardScaler
from sklearn.metrics import mean_absolute_error, mean_squared_error, r2_score
from sklearn.metrics import accuracy_score, precision_score, recall_score, f1_score, confusion_matrix

from sklearn.linear_model import LinearRegression
from sklearn.ensemble import RandomForestRegressor
from xgboost import XGBRegressor

# Nëse do LSTM/GRU
from tensorflow.keras.models import Sequential
from tensorflow.keras.layers import LSTM, GRU, Dense
from tensorflow.keras.callbacks import EarlyStopping

USE_SQLITE = False

CSV_PATH = r"D:\Doktoratura\Air quality data\viti2025+.csv"
SQLITE_PATH = "digital_twin.db"
TABLE_NAME = "sensor_data"

timestamp_col = "Datetime"
humidity_col = "Lageshtira"
```

```

if USE_SQLITE:
    conn = sqlite3.connect(SQLITE_PATH)
    df = pd.read_sql_query(f"SELECT * FROM {TABLE_NAME}", conn)
    conn.close()
else:
    df = pd.read_csv(CSV_PATH)
df[timestamp_col] = pd.to_datetime(df[timestamp_col], errors="coerce")
df[humidity_col] = pd.to_numeric(df[humidity_col], errors="coerce")

df = df.dropna(subset=[timestamp_col, humidity_col])
df = df.sort_values(timestamp_col).reset_index(drop=True)

# =====
# HORIZONTI 30 MINUTA
# Supozojmë matje çdo 5 minuta => 30/5 = 6 hapa
# =====

sampling_minutes = 5
horizon_minutes = 30
steps_ahead = int(horizon_minutes / sampling_minutes)

df["target_humidity_30min"] = df[humidity_col].shift(-steps_ahead)

# =====
# VEÇORITË
# =====

df["humidity_lag_1"] = df[humidity_col].shift(1)
df["humidity_lag_2"] = df[humidity_col].shift(2)
df["humidity_lag_3"] = df[humidity_col].shift(3)
df["humidity_lag_6"] = df[humidity_col].shift(6)
df["humidity_lag_12"] = df[humidity_col].shift(12)

df["rolling_mean_6"] = df[humidity_col].rolling(window=6).mean()
df["rolling_std_6"] = df[humidity_col].rolling(window=6).std()

df["hour"] = df[timestamp_col].dt.hour
df["dayofweek"] = df[timestamp_col].dt.dayofweek
df["month"] = df[timestamp_col].dt.month

df = df.dropna().reset_index(drop=True)
features = [
    humidity_col,
    "humidity_lag_1",
    "humidity_lag_2",
    "humidity_lag_3",
    "humidity_lag_6",
    "humidity_lag_12",
    "rolling_mean_6",
    "rolling_std_6",
    "hour",
    "dayofweek",
    "month"
]

target = "target_humidity_30min"

```

```

X = df[features]
y = df[target]

# =====
# NDARJA KRONOLOGJIKE TRAIN / TEST
# =====

split_index = int(len(df) * 0.8)

X_train = X.iloc[:split_index]
X_test = X.iloc[split_index:]

y_train = y.iloc[:split_index]
y_test = y.iloc[split_index:]

# =====
# MODELET ML REGRESIVE
# =====

models = {
    "Linear Regression": LinearRegression(),
    "Random Forest": RandomForestRegressor(
        n_estimators=100,
        random_state=42,
        n_jobs=-1
    ),
    "XGBoost": XGBRegressor(
        n_estimators=150,
        learning_rate=0.05,
        max_depth=5,
        random_state=42,
        objective="reg:squarederror"
    )
}

results = []

THRESHOLD = 69

def evaluate_model(model_name, y_true, y_pred):
    mae = mean_absolute_error(y_true, y_pred)
    rmse = np.sqrt(mean_squared_error(y_true, y_pred))
    r2 = r2_score(y_true, y_pred)

    # Vendim me prag të vetëm
    # 1 = lagështira është mbi ose baraz me prapun
    # 0 = lagështira është nën prag
    y_true_class = (y_true >= THRESHOLD).astype(int)
    y_pred_class = (y_pred >= THRESHOLD).astype(int)

    acc = accuracy_score(y_true_class, y_pred_class)
    precision = precision_score(y_true_class, y_pred_class, zero_division=0)

```

```

recall = recall_score(y_true_class, y_pred_class, zero_division=0)
f1 = f1_score(y_true_class, y_pred_class, zero_division=0)

tn, fp, fn, tp = confusion_matrix(y_true_class, y_pred_class).ravel()

results.append({
    "Model": model_name,
    "MAE": mae,
    "RMSE": rmse,
    "R2": r2,
    "Accuracy": acc,
    "Precision": precision,
    "Recall": recall,
    "F1": f1,
    "TP": tp,
    "TN": tn,
    "FP": fp,
    "FN": fn
})

for model_name, model in models.items():
    print(f"Trajnimi i modelit: {model_name}")
    model.fit(X_train, y_train)
    y_pred = model.predict(X_test)
    evaluate_model(model_name, y_test, y_pred)

# =====
# LSTM DHE GRU modelet DL
# =====

scaler_X = StandardScaler()
scaler_y = StandardScaler()

X_train_scaled = scaler_X.fit_transform(X_train)
X_test_scaled = scaler_X.transform(X_test)

y_train_scaled = scaler_y.fit_transform(y_train.values.reshape(-1, 1))
y_test_scaled = scaler_y.transform(y_test.values.reshape(-1, 1))

def create_sequences(X_data, y_data, seq_len=12):
    X_seq, y_seq = [], []
    for i in range(seq_len, len(X_data)):
        X_seq.append(X_data[i-seq_len:i])
        y_seq.append(y_data[i])
    return np.array(X_seq), np.array(y_seq)

seq_len = 12 # 12 hapa nga 5 minuta = 60 minuta histori

X_train_seq, y_train_seq = create_sequences(X_train_scaled, y_train_scaled, seq_len)
X_test_seq, y_test_seq = create_sequences(X_test_scaled, y_test_scaled, seq_len)

y_test_real_seq = scaler_y.inverse_transform(y_test_seq).flatten()

```

```

def build_lstm(input_shape):
    model = Sequential()
    model.add(LSTM(64, input_shape=input_shape))
    model.add(Dense(1))
    model.compile(optimizer="adam", loss="mse")
    return model

def build_gru(input_shape):
    model = Sequential()
    model.add(GRU(64, input_shape=input_shape))
    model.add(Dense(1))
    model.compile(optimizer="adam", loss="mse")
    return model

early_stop = EarlyStopping(
    monitor="val_loss",
    patience=5,
    restore_best_weights=True
)

deep_models = {
    "LSTM": build_lstm((X_train_seq.shape[1], X_train_seq.shape[2])),
    "GRU": build_gru((X_train_seq.shape[1], X_train_seq.shape[2]))
}

for model_name, model in deep_models.items():
    print(f"Trajnimi i modelit: {model_name}")

    model.fit(
        X_train_seq,
        y_train_seq,
        epochs=30,
        batch_size=64,
        validation_split=0.2,
        callbacks=[early_stop],
        verbose=1
    )

    y_pred_scaled = model.predict(X_test_seq)
    y_pred_real = scaler_y.inverse_transform(y_pred_scaled).flatten()

    evaluate_model(model_name, y_test_real_seq, y_pred_real)

# =====
# TABELA FINALE
# =====

results_df = pd.DataFrame(results)

```

```

results_df = results_df.sort_values(
    by=["RMSE", "MAE", "F1"],
    ascending=[True, True, False]
)

print("\nRezultatet finale:")
print(results_df)

results_df.to_csv("model_comparison_30min_pragu_69.csv", index=False)

print("\nModeli më i mirë sipas RMSE:")
print(results_df.iloc[0])

```

Bibliografia

- [1] A. M. Turing, "I.—COMPUTING MACHINERY AND INTELLIGENCE," *Mind*, vol. LIX, no. 236, pp. 433–460, Oct. 1950, doi: 10.1093/mind/LIX.236.433.
- [2] M. Haenlein and A. Kaplan, "A Brief History of Artificial Intelligence: On the Past, Present, and Future of Artificial Intelligence," *California Management Review*, vol. 61, no. 4, pp. 5–14, Aug. 2019, doi: 10.1177/0008125619864925.
- [3] Y. Xu *et al.*, "Artificial intelligence: A powerful paradigm for scientific research," *The Innovation*, vol. 2, no. 4, p. 100179, Nov. 2021, doi: 10.1016/j.xinn.2021.100179.
- [4] S. Makridakis, "The forthcoming Artificial Intelligence (AI) revolution: Its impact on society and firms," *Futures*, vol. 90, pp. 46–60, Jun. 2017, doi: 10.1016/j.futures.2017.03.006.
- [5] E. Brynjolfsson and A. McAfee, *The second machine age: work, progress, and prosperity in a time of brilliant technologies*, First published as a Norton paperback. New York London: W. W. Norton & Company, 2016.
- [6] Y. R. Shrestha, S. M. Ben-Menahem, and G. Von Krogh, "Organizational Decision-Making Structures in the Age of Artificial Intelligence," *California Management Review*, vol. 61, no. 4, pp. 66–83, Aug. 2019, doi: 10.1177/0008125619862257.
- [7] M. H. Jarrahi, "Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making," *Business Horizons*, vol. 61, no. 4, pp. 577–586, Jul. 2018, doi: 10.1016/j.bushor.2018.03.007.
- [8] V. Mayer-Schönberger and K. Cukier, *Big data: a revolution that will transform how we live, work, and think*. Boston: Houghton Mifflin Harcourt, 2013.
- [9] R. Kitchin, *The Data Revolution: Big Data, Open Data, Data Infrastructures & Their Consequences*. 1 Oliver's Yard, 55 City Road, London EC1Y 1SP United Kingdom: SAGE Publications Ltd, 2014. doi: 10.4135/9781473909472.
- [10] E. Brynjolfsson and A. McAfee, *The second machine age: work, progress, and prosperity in a time of brilliant technologies*, First published as a Norton paperback. New York London: W. W. Norton & Company, 2016.
- [11] R. Parasuraman, T. B. Sheridan, and C. D. Wickens, "A model for types and levels of human interaction with automation," *IEEE Trans. Syst., Man, Cybern. A*, vol. 30, no. 3, pp. 286–297, May 2000, doi: 10.1109/3468.844354.
- [12] D. Castelvechi, "Can we open the black box of AI?," *Nature*, vol. 538, no. 7623, pp. 20–23, Oct. 2016, doi: 10.1038/538020a.
- [13] M. Stonebraker, D. Bruckner, I. F. Ilyas, A. Pagan, and S. Xu, "Data Curation at Scale: The Data Tamer System," 2013.
- [14] P. Joshi, S. Santy, A. Budhiraja, K. Bali, and M. Choudhury, "The State and Fate of Linguistic Diversity and Inclusion in the NLP World," in *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, Online: Association for Computational Linguistics, 2020, pp. 6282–6293. doi: 10.18653/v1/2020.acl-main.560.
- [15] D. I. Adelani *et al.*, "MasakhaNER: Named Entity Recognition for African Languages," *Transactions of the Association for Computational Linguistics*, vol. 9, pp. 1116–1131, Oct. 2021, doi: 10.1162/tacl_a_00416.
- [16] T. H. Davenport and R. RONANKI, "Artificial Intelligence for the Real World," *HARVARD BUSINESS REVIEW*, 2018.
- [17] S. Ivanov and C. Webster, "Conceptual Framework of the Use of Robots, Artificial Intelligence and Service Automation in Travel, Tourism, and Hospitality Companies," in *Robots, Artificial Intelligence, and Service Automation in Travel, Tourism and Hospitality*, S. Ivanov and C. Webster, Eds., Emerald Publishing Limited, 2019, pp. 7–37. doi: 10.1108/978-1-78756-687-320191001.

- [18] G. Piras, S. Agostinelli, and F. Muzi, "Smart Buildings and Digital Twin to Monitoring the Efficiency and Wellness of Working Environments: A Case Study on IoT Integration and Data-Driven Management," *Applied Sciences*, vol. 15, no. 9, p. 4939, Apr. 2025, doi: 10.3390/app15094939.
- [19] M. I. Jordan and T. M. Mitchell, "Machine learning: Trends, perspectives, and prospects," *Science*, vol. 349, no. 6245, pp. 255–260, Jul. 2015, doi: 10.1126/science.aaa8415.
- [20] S. J. Russell and P. Norvig, *Artificial intelligence: a modern approach*, Fourth edition, Global edition. in Prentice Hall series in artificial intelligence. Boston: Pearson, 2022.
- [21] R. Kitchin, "Thinking critically about and researching algorithms," *Information, Communication & Society*, vol. 20, no. 1, pp. 14–29, Jan. 2017, doi: 10.1080/1369118X.2016.1154087.
- [22] F. Doshi-Velez and B. Kim, "Towards A Rigorous Science of Interpretable Machine Learning," 2017, *arXiv*. doi: 10.48550/ARXIV.1702.08608.
- [23] R. Guidotti, A. Monreale, S. Ruggieri, F. Turini, D. Pedreschi, and F. Giannotti, "A Survey Of Methods For Explaining Black Box Models," 2018, *arXiv*. doi: 10.48550/ARXIV.1802.01933.
- [24] L. Floridi *et al.*, "AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations," *Minds & Machines*, vol. 28, no. 4, pp. 689–707, Dec. 2018, doi: 10.1007/s11023-018-9482-5.
- [25] S. Wachter, B. Mittelstadt, and C. Russell, "Counterfactual Explanations without Opening the Black Box: Automated Decisions and the GDPR," 2017, doi: 10.48550/ARXIV.1711.00399.
- [26] Y. Duan, J. S. Edwards, and Y. K. Dwivedi, "Artificial intelligence for decision making in the era of Big Data – evolution, challenges and research agenda," *International Journal of Information Management*, vol. 48, pp. 63–71, Oct. 2019, doi: 10.1016/j.ijinfomgt.2019.01.021.
- [27] M. Ravi, A. Negi, N. S. Bommi, and N. Rouf, "Evolution of AI-Driven Decision Making with Decision Support Systems, Expert Systems, Recommender Systems, and XAI," *IETE Technical Review*, vol. 42, no. 4, pp. 428–465, Jul. 2025, doi: 10.1080/02564602.2025.2512086.
- [28] A. Tubman, "The Use of Artificial Intelligence in International Decision-Making Processes in Project Management," *SSRN Journal*, 2022, doi: 10.2139/ssrn.4121200.
- [29] R. Sultana, "Artificial Intelligence For Decision Making In The Era Of Big Data Evolution," *NHJ*, vol. 1, no. 01, pp. 17–40, Nov. 2024, doi: 10.70008/jbimistr.v1i01.59.
- [30] E. Yalunina, N. Pryadilina, and E. Skvorcov, "Improving the process of making management decisions in agriculture using artificial intelligence systems," *Agrarian Bulletin of the*, vol. 24, no. 03, pp. 440–449, Mar. 2024, doi: 10.32417/1997-4868-2024-24-03-440-449.
- [31] T. Davenport and R. Kalakota, "The potential for artificial intelligence in healthcare," *Future Healthcare Journal*, vol. 6, no. 2, pp. 94–98, Jun. 2019, doi: 10.7861/futurehosp.6-2-94.
- [32] E. J. Topol, "High-performance medicine: the convergence of human and artificial intelligence," *Nat Med*, vol. 25, no. 1, pp. 44–56, Jan. 2019, doi: 10.1038/s41591-018-0300-7.
- [33] A. Barredo Arrieta *et al.*, "Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI," *Information Fusion*, vol. 58, pp. 82–115, Jun. 2020, doi: 10.1016/j.inffus.2019.12.012.
- [34] V. Lai, C. Chen, Q. V. Liao, A. Smith-Renner, and C. Tan, "Towards a Science of Human-AI Decision Making: A Survey of Empirical Studies," 2021, *arXiv*. doi: 10.48550/ARXIV.2112.11471.
- [35] J. J. Ferreira and M. Monteiro, "Designer-User Communication for XAI: An

- epistemological approach to discuss XAI design,” 2021, *arXiv*. doi: 10.48550/ARXIV.2105.07804.
- [36] A. K. Boos, “Conceptualizing Automated Decision-Making in Organizational Contexts,” *Philos. Technol.*, vol. 37, no. 3, p. 92, Sep. 2024, doi: 10.1007/s13347-024-00773-5.
 - [37] G. De Gasperis and S. D. Facchini, “A Comparative Study of Rule-Based and Data-Driven Approaches in Industrial Monitoring,” 2025, *arXiv*. doi: 10.48550/ARXIV.2509.15848.
 - [38] S. J. Russell and P. Norvig, *Artificial intelligence: a modern approach*, Fourth edition, Global edition. in Prentice Hall series in artificial intelligence. Boston: Pearson, 2022.
 - [39] R. Parasuraman and V. Riley, “Humans and Automation: Use, Misuse, Disuse, Abuse,” *Hum Factors*, vol. 39, no. 2, pp. 230–253, Jun. 1997, doi: 10.1518/001872097778543886.
 - [40] N. Diakopoulos, “Accountability in algorithmic decision making,” *Commun. ACM*, vol. 59, no. 2, pp. 56–62, Jan. 2016, doi: 10.1145/2844110.
 - [41] P. Resnick and H. R. Varian, “Recommender systems,” *Commun. ACM*, vol. 40, no. 3, pp. 56–58, Mar. 1997, doi: 10.1145/245108.245121.
 - [42] P. Rajpurkar, E. Chen, O. Banerjee, and E. J. Topol, “AI in health and medicine,” *Nat Med*, vol. 28, no. 1, pp. 31–38, Jan. 2022, doi: 10.1038/s41591-021-01614-0.
 - [43] E. J. Topol and A. Verghese, *Deep medicine: how artificial intelligence can make healthcare human again*, First edition. New York, NY: Basic Books, 2019.
 - [44] T. Hendershott and R. Riordan, “Algorithmic Trading and the Market for Liquidity,” *J. Financ. Quant. Anal.*, vol. 48, no. 4, pp. 1001–1024, Aug. 2013, doi: 10.1017/S0022109013000471.
 - [45] S. G. Pauker, G. A. Gorry, J. P. Kassirer, and W. B. Schwartz, “Towards the simulation of clinical cognition,” *The American Journal of Medicine*, vol. 60, no. 7, pp. 981–996, Jun. 1976, doi: 10.1016/0002-9343(76)90570-2.
 - [46] T. M. Mitchell, *Machine learning*, Nachdr. in McGraw-Hill series in Computer Science. New York: McGraw-Hill, 1997.
 - [47] F. Provost and T. Fawcett, *Data science for business: what you need to know about data mining and data-analytic thinking*, First edition. Beijing Cambridge Farnham Köln Sebastopol Tokyo: O’Reilly, 2013.
 - [48] T. B. Brown *et al.*, “Language Models are Few-Shot Learners,” 2020, *arXiv*. doi: 10.48550/ARXIV.2005.14165.
 - [49] Y. LeCun, Y. Bengio, and G. Hinton, “Deep learning,” *Nature*, vol. 521, no. 7553, pp. 436–444, May 2015, doi: 10.1038/nature14539.
 - [50] A. Vaswani *et al.*, “Attention Is All You Need,” 2017, *arXiv*. doi: 10.48550/ARXIV.1706.03762.
 - [51] S. Wachter, B. Mittelstadt, and L. Floridi, “Why a Right to Explanation of Automated Decision-Making Does Not Exist in the General Data Protection Regulation,” *International Data Privacy Law*, vol. 7, no. 2, pp. 76–99, May 2017, doi: 10.1093/idpl/ix005.
 - [52] L. Floridi *et al.*, “AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations,” *Minds & Machines*, vol. 28, no. 4, pp. 689–707, Dec. 2018, doi: 10.1007/s11023-018-9482-5.
 - [53] R. Parasuraman, T. B. Sheridan, and C. D. Wickens, “A model for types and levels of human interaction with automation,” *IEEE Trans. Syst., Man, Cybern. A*, vol. 30, no. 3, pp. 286–297, May 2000, doi: 10.1109/3468.844354.
 - [54] Y. R. Shrestha, S. M. Ben-Menahem, and G. Von Krogh, “Organizational Decision-Making Structures in the Age of Artificial Intelligence,” *California Management Review*, vol. 61, no. 4, pp. 66–83, Aug. 2019, doi: 10.1177/0008125619862257.
 - [55] R. Binns, M. Van Kleek, M. Veale, U. Lyngs, J. Zhao, and N. Shadbolt, “‘It’s Reducing a Human Being to a Percentage’: Perceptions of Justice in Algorithmic Decisions,” in *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*, Montreal QC Canada: ACM, Apr. 2018, pp. 1–14. doi: 10.1145/3173574.3173951.

- [56] H. A. Simon, “A Behavioral Model of Rational Choice,” *The Quarterly Journal of Economics*, vol. 69, no. 1, p. 99, Feb. 1955, doi: 10.2307/1884852.
- [57] D. Kahneman and A. Tversky, “Prospect Theory: An Analysis of Decision under Risk,” *Econometrica*, vol. 47, no. 2, p. 263, Mar. 1979, doi: 10.2307/1914185.
- [58] J. Von Neumann and O. Morgenstern, *Theory of Games and Economic Behavior (60th Anniversary Commemorative Edition)*: Princeton University Press, 2007. doi: 10.1515/9781400829460.
- [59] D. Koller and N. Friedman, *Probabilistic graphical models: principles and techniques*. in Adaptive computation and machine learning. Cambridge, MA: MIT Press, 2009.
- [60] J. Von Neumann and O. Morgenstern, *Theory of games and economic behavior*, 60. anniversary ed., 4. print., and 1. paperb. print. in A Princeton classic edition. Princeton, NJ: Princeton University Press, 2007.
- [61] G. James, D. Witten, T. Hastie, and R. Tibshirani, *An Introduction to Statistical Learning*, vol. 103. in Springer Texts in Statistics, vol. 103. New York, NY: Springer New York, 2013. doi: 10.1007/978-1-4614-7138-7.
- [62] L. Breiman, J. H. Friedman, R. A. Olshen, and C. J. Stone, *Classification And Regression Trees*, 1st ed. Routledge, 2017. doi: 10.1201/9781315139470.
- [63] L. Breiman, “Random Forests,” *Machine Learning*, vol. 45, no. 1, pp. 5–32, Oct. 2001, doi: 10.1023/A:1010933404324.
- [64] T. Chen and C. Guestrin, “XGBoost: A Scalable Tree Boosting System,” in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, San Francisco California USA: ACM, Aug. 2016, pp. 785–794. doi: 10.1145/2939672.2939785.
- [65] J. H. Friedman, “Greedy function approximation: A gradient boosting machine.,” *Ann. Statist.*, vol. 29, no. 5, Oct. 2001, doi: 10.1214/aos/1013203451.
- [66] C. Cortes and V. Vapnik, “Support-vector networks,” *Mach Learn*, vol. 20, no. 3, pp. 273–297, Sep. 1995, doi: 10.1007/BF00994018.
- [67] S. Hochreiter and J. Schmidhuber, “Long Short-Term Memory,” *Neural Computation*, vol. 9, no. 8, pp. 1735–1780, Nov. 1997, doi: 10.1162/neco.1997.9.8.1735.
- [68] D. Bahdanau, K. Cho, and Y. Bengio, “Neural Machine Translation by Jointly Learning to Align and Translate,” 2014, *arXiv*. doi: 10.48550/ARXIV.1409.0473.
- [69] B. W. Wirtz, J. C. Weyerer, and C. Geyer, “Artificial Intelligence and the Public Sector—Applications and Challenges,” *International Journal of Public Administration*, vol. 42, no. 7, pp. 596–615, May 2019, doi: 10.1080/01900692.2018.1498103.
- [70] J. Amann, A. Blasimme, E. Vayena, D. Frey, and V. I. Madai, “Explainability for artificial intelligence in healthcare: a multidisciplinary perspective,” *BMC Med Inform Decis Mak*, vol. 20, no. 1, p. 310, Dec. 2020, doi: 10.1186/s12911-020-01332-6.
- [71] W.-N. Hsu, B. Bolte, Y.-H. H. Tsai, K. Lakhotia, R. Salakhutdinov, and A. Mohamed, “HuBERT: Self-Supervised Speech Representation Learning by Masked Prediction of Hidden Units,” 2021, *arXiv*. doi: 10.48550/ARXIV.2106.07447.
- [72] D. B. Resnik and M. Hosseini, “The ethics of using artificial intelligence in scientific research: new guidance needed for a new tool,” *AI Ethics*, vol. 5, no. 2, pp. 1499–1521, Apr. 2025, doi: 10.1007/s43681-024-00493-8.
- [73] M. Ridley, “Explainable Artificial Intelligence (XAI): Adoption and Advocacy,” *ITAL*, vol. 41, no. 2, Jun. 2022, doi: 10.6017/ital.v41i2.14683.
- [74] S. Lundberg and S.-I. Lee, “A Unified Approach to Interpreting Model Predictions,” 2017, *arXiv*. doi: 10.48550/ARXIV.1705.07874.
- [75] M. T. Ribeiro, S. Singh, and C. Guestrin, ““Why Should I Trust You?”: Explaining the Predictions of Any Classifier,” in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, San Francisco California USA: ACM, Aug. 2016, pp. 1135–1144. doi: 10.1145/2939672.2939778.

- [76] R. Parasuraman, T. B. Sheridan, and C. D. Wickens, “A model for types and levels of human interaction with automation,” *IEEE Trans. Syst., Man, Cybern. A*, vol. 30, no. 3, pp. 286–297, May 2000, doi: 10.1109/3468.844354.
- [77] S. Boyd and L. Vandenberghe, *Convex Optimization*, 1st ed. Cambridge University Press, 2004. doi: 10.1017/CBO9780511804441.
- [78] R. Kitchin, “Thinking critically about and researching algorithms,” *Information, Communication & Society*, vol. 20, no. 1, pp. 14–29, Jan. 2017, doi: 10.1080/1369118X.2016.1154087.
- [79] F. Provost and T. Fawcett, *Data science for business: what you need to know about data mining and data-analytic thinking*, First edition. Beijing Cambridge Farnkam Köln Sebastopol Tokyo: O’Reilly, 2013.
- [80] I. Goodfellow, Y. Bengio, and A. Courville, *Deep learning*. in Adaptive computation and machine learning. Cambridge, Massachusetts: The MIT Press, 2016.
- [81] S. J. Russell and P. Norvig, *Artificial intelligence: a modern approach*, Fourth Edition. in Pearson Series in Artificial Intelligence. Hoboken, NJ: Pearson, 2021.
- [82] Y. LeCun, Y. Bengio, and G. Hinton, “Deep learning,” *Nature*, vol. 521, no. 7553, pp. 436–444, May 2015, doi: 10.1038/nature14539.
- [83] G. Shmueli and O. R. Koppius, “Predictive analytics in information systems research,” vol. 35, no. 3, OHDSI, pp. 553–572, 2011. doi: 10.1287/isre.1100.0347.
- [84] F. Eyben *et al.*, “The Geneva Minimalistic Acoustic Parameter Set (GeMAPS) for Voice Research and Affective Computing,” *IEEE Trans. Affective Comput.*, vol. 7, no. 2, pp. 190–202, Apr. 2016, doi: 10.1109/TAFFC.2015.2457417.
- [85] D. Jurafsky and J. H. Martin, *Speech and language processing: an introduction to natural language processing, computational linguistics, and speech recognition*, 2. ed. [Nachdr.]. Upper Saddle River, NJ: Prentice Hall, 2009.
- [86] A. Baeveski, H. Zhou, A. Mohamed, and M. Auli, “wav2vec 2.0: A Framework for Self-Supervised Learning of Speech Representations,” 2020, *arXiv*. doi: 10.48550/ARXIV.2006.11477.
- [87] D. Jurafsky and J. H. Martin, *Speech and language processing: an introduction to natural language processing, computational linguistics, and speech recognition*, 2nd ed. in Prentice Hall series in artificial intelligence. Upper Saddle River, N.J: Pearson Prentice Hall, 2009.
- [88] P. C. Loizou, *Speech Enhancement: Theory and Practice*, 2nd ed. CRC Press, 2013. doi: 10.1201/b14529.
- [89] L. R. Rabiner and R. W. Schafer, *Introduction to digital speech processing*. in Foundations and trends in signal processing, no. v. 1, Issue 1-2, 2007. Boston, Mass: Now, 2007.
- [90] L. R. Rabiner and R. W. Schafer, “Introduction to Digital Speech Processing,” *Foundations and Trends® in Signal Processing*, vol. 1, no. 1–2, pp. 1–194, Dec. 2007, doi: 10.1561/20000000001.
- [91] C. M. Bishop, *Pattern Recognition and Machine Learning*. Springer New York, 2006. doi: 10.1007/978-0-387-45528-0.
- [92] B. Schuller, A. Batliner, S. Steidl, and D. Seppi, “Recognising realistic emotions and affect in speech: State of the art and lessons learnt from the first challenge,” *Speech Communication*, vol. 53, no. 9–10, pp. 1062–1087, Nov. 2011, doi: 10.1016/j.specom.2011.01.011.
- [93] A. Vaswani *et al.*, “Attention Is All You Need,” 2017, *arXiv*. doi: 10.48550/ARXIV.1706.03762.
- [94] I. T. Jolliffe, *Principal Component Analysis*. in Springer Series in Statistics. New York: Springer-Verlag, 2002. doi: 10.1007/b98835.
- [95] M. Sokolova and G. Lapalme, “A systematic analysis of performance measures for classification tasks,” *Information Processing & Management*, vol. 45, no. 4, pp. 427–437, Jul. 2009, doi: 10.1016/j.ipm.2009.03.002.

- [96] L. van der Maaten and G. Hinton, "Visualizing Data using t-SNE," *Journal of Machine Learning Research* 9, vol. 9, no. 2579–2605, 2008.
- [97] S. Batra and A. Chatterji, "Enhancing Customer Experience in the Hospitality Industry through Artificial Intelligence," *Don Bosco Institute of Technology Delhi Journal of Research*, vol. 1, no. 1, pp. 6–12, Sep. 2024, doi: 10.48165/dbitdj.2024.1.01.02.
- [98] C. Prentice, S. Dominique Lopes, and X. Wang, "The impact of artificial intelligence and employee service quality on customer satisfaction and loyalty," *Journal of Hospitality Marketing & Management*, vol. 29, no. 7, pp. 739–756, Oct. 2020, doi: 10.1080/19368623.2020.1722304.
- [99] H. S. Al-Hyari and H. M. Al-Smadi, "The Impact of Artificial Intelligence (ai) on Guest Satisfaction in Hotel Management: An Empirical Study of Luxury Hotels," *GTG*, vol. 48, no. 2 supplement, pp. 810–819, Jun. 2023, doi: 10.30892/gtg.482spl15-1081.
- [100] A. Rawat, R. Kukreti, A. Dimari, and R. Dani, "Artificial Intelligence in HMI System," in *2024 4th International Conference on Advance Computing and Innovative Technologies in Engineering (ICACITE)*, Greater Noida, India: IEEE, May 2024, pp. 1362–1367. doi: 10.1109/ICACITE60783.2024.10617209.
- [101] A. Kosta, F. Pazari, I. Lili, and S. Greca, "The Role of Artificial Intelligence in Transforming Tourism Services: The Case of Durrës, Albania," *GTG*, vol. 61, no. 3, pp. 1485–1494, Sep. 2025, doi: 10.30892/gtg.61308-1518.
- [102] F. Huseynov and Y. C. Güler, "The Impact of Online Consumer Reviews on Online Hotel Booking Intention," *isarder*, vol. 13, no. 3, pp. 2634–2652, Sep. 2021, doi: 10.20491/isarder.2021.1282.
- [103] J. Amann, A. Blasimme, E. Vayena, D. Frey, and V. I. Madai, "Explainability for artificial intelligence in healthcare: a multidisciplinary perspective," *BMC Med Inform Decis Mak*, vol. 20, no. 1, p. 310, Dec. 2020, doi: 10.1186/s12911-020-01332-6.
- [104] S. Ali *et al.*, "Explainable Artificial Intelligence (XAI): What we know and what is left to attain Trustworthy Artificial Intelligence," *Information Fusion*, vol. 99, p. 101805, Nov. 2023, doi: 10.1016/j.inffus.2023.101805.
- [105] R. Tiwari, "Explainable AI (XAI) and its Applications in Building Trust and Understanding in AI Decision Making," *IJSREM*, vol. 07, no. 01, Jan. 2023, doi: 10.55041/IJSREM17592.
- [106] J. Singh, S. Rani, and G. Srilakshmi, "Towards Explainable AI: Interpretable Models for Complex Decision-making," in *2024 International Conference on Knowledge Engineering and Communication Systems (ICKECS)*, Chikkaballapur, India: IEEE, Apr. 2024, pp. 1–5. doi: 10.1109/ICKECS61492.2024.10616500.
- [107] G. P. Reddy and Y. V. P. Kumar, "Explainable AI (XAI): Explained," in *2023 IEEE Open Conference of Electrical, Electronic and Information Sciences (eStream)*, Vilnius, Lithuania: IEEE, Apr. 2023, pp. 1–6. doi: 10.1109/eStream59056.2023.10134984.
- [108] M. Bhuiyan, M. A. Rahman, A. B. Hoque, M. Ashrafuzzaman, and A. Rahman, "Advanced Analytics and Machine Learning for Revenue Optimization in the Hospitality Industry: A Comprehensive Review of Frameworks," *AJSRI*, vol. 2, no. 02, pp. 52–74, Sep. 2023, doi: 10.63125/8xbkma40.
- [109] F. A. Anwar, D. Deliana, and S. Suyamto, "Digital Transformation in the Hospitality Industry: Improving Efficiency and Guest Experience," *IJMSIT*, vol. 4, no. 2, pp. 428–437, Sep. 2024, doi: 10.35870/ijmsit.v4i2.3201.
- [110] M. M. Ridzky, "Studies Effectiveness System Management Operational For Optimization Costs in Industry Hospitality," *hmoj*, vol. 1, no. 1, pp. 36–45, Dec. 2024, doi: 10.59261/hmoj.v1i1.5.
- [111] M. Castelli, D. C. Pinto, S. Shuqair, D. Montali, and L. Vanneschi, "The Benefits of Automated Machine Learning in Hospitality: A Step-By-Step Guide and AutoML Tool," *Emerg Sci J*, vol. 6, no. 6, pp. 1237–1254, Sep. 2022, doi: 10.28991/ESJ-2022-06-06-02.
- [112] A. Das and P. Rad, "Opportunities and Challenges in Explainable Artificial Intelligence

- (XAI): A Survey,” 2020, *arXiv*. doi: 10.48550/ARXIV.2006.11371.
- [113] T. A. Roshinta and S. Gábor, “A Comparative Study of LIME and SHAP for Enhancing Trustworthiness and Efficiency in Explainable AI Systems,” in *2024 IEEE International Conference on Computing (ICOCO)*, Kuala Lumpur, Malaysia: IEEE, Dec. 2024, pp. 134–139. doi: 10.1109/ICOCO62848.2024.10928183.
 - [114] C. R. Jones and B. K. Bergen, “Large Language Models Pass the Turing Test,” 2025, *arXiv*. doi: 10.48550/ARXIV.2503.23674.
 - [115] A. Rahimov, O. Zamler, and A. Azaria, “The Turing Test Is More Relevant Than Ever,” 2025, *arXiv*. doi: 10.48550/ARXIV.2505.02558.
 - [116] A. Gupta, “Introduction to AI Chatbots,” *IJERT*, vol. V9, no. 07, p. IJERTV9IS070143, Jul. 2020, doi: 10.17577/IJERTV9IS070143.
 - [117] S. Hussain, O. Ameri Sianaki, and N. Ababneh, “A Survey on Conversational Agents/Chatbots Classification and Design Techniques,” in *Web, Artificial Intelligence and Network Applications*, vol. 927, L. Barolli, M. Takizawa, F. Xhafa, and T. Enokido, Eds., in *Advances in Intelligent Systems and Computing*, vol. 927., Cham: Springer International Publishing, 2019, pp. 946–956. doi: 10.1007/978-3-030-15035-8_93.
 - [118] Y. E. Madhi and B. S. Mustafa, “History of Chatbot from Past to Present,” *IRJIET*, vol. 08, no. 11, pp. 01–05, 2024, doi: 10.47001/IRJIET/2024.811001.
 - [119] K. Wang, “From ELIZA to ChatGPT: A brief history of chatbots and their evolution,” *ACE*, vol. 39, no. 1, pp. 57–62, Feb. 2024, doi: 10.54254/2755-2721/39/20230579.
 - [120] Md. Al-Amin *et al.*, “History of generative Artificial Intelligence (AI) chatbots: past, present, and future development,” 2024, *arXiv*. doi: 10.48550/ARXIV.2402.05122.
 - [121] D. Buhalis, P. O’Connor, and R. Leung, “Smart hospitality: from smart cities and smart tourism towards agile business ecosystems in networked destinations,” *IJCHM*, vol. 35, no. 1, pp. 369–393, Jan. 2023, doi: 10.1108/IJCHM-04-2022-0497.
 - [122] L. Mich and R. Garigliano, “ChatGPT for e-Tourism: a technological perspective,” *Inf Technol Tourism*, vol. 25, no. 1, pp. 1–12, Mar. 2023, doi: 10.1007/s40558-023-00248-x.
 - [123] A. M. Fouad, I. E. Salem, and E. A. Fathy, “Generative AI insights in tourism and hospitality: A comprehensive review and strategic research roadmap,” *Tourism and Hospitality Research*, p. 14673584241293125, Oct. 2024, doi: 10.1177/14673584241293125.
 - [124] D. Naik, I. Naik, and N. Naik, “Applications of AI Chatbots Based on Generative AI, Large Language Models and Large Multimodal Models,” in *Contributions Presented at The International Conference on Computing, Communication, Cybersecurity and AI, July 3–4, 2024, London, UK*, vol. 884, N. Naik, P. Jenkins, S. Prajapat, and P. Grace, Eds., in *Lecture Notes in Networks and Systems*, vol. 884., Cham: Springer Nature Switzerland, 2024, pp. 668–690. doi: 10.1007/978-3-031-74443-3_39.
 - [125] D. Naik, I. Naik, and N. Naik, “Applications of AI Chatbots Based on Generative AI, Large Language Models and Large Multimodal Models,” Nov. 10, 2024. doi: 10.36227/techrxiv.173121409.91676949/v1.
 - [126] H. Touvron *et al.*, “LLaMA: Open and Efficient Foundation Language Models,” 2023, *arXiv*. doi: 10.48550/ARXIV.2302.13971.
 - [127] S. Ali *et al.*, “Explainable Artificial Intelligence (XAI): What we know and what is left to attain Trustworthy Artificial Intelligence,” *Information Fusion*, vol. 99, p. 101805, Nov. 2023, doi: 10.1016/j.inffus.2023.101805.
 - [128] T. Zhang, M. Zhang, W. Y. Low, X. J. Yang, and B. A. Li, “Conversational Explanations: Discussing Explainable AI with Non-AI Experts,” in *Proceedings of the 30th International Conference on Intelligent User Interfaces*, Cagliari Italy: ACM, Mar. 2025, pp. 409–424. doi: 10.1145/3708359.3712143.
 - [129] G. He, N. Aishwarya, and U. Gadiraju, “Is Conversational XAI All You Need? Human-AI Decision Making With a Conversational XAI Assistant,” in *Proceedings of the 30th*

- International Conference on Intelligent User Interfaces*, Cagliari Italy: ACM, Mar. 2025, pp. 907–924. doi: 10.1145/3708359.3712133.
- [130] M. Ridley, “Explainable Artificial Intelligence (XAI): Adoption and Advocacy,” *ITAL*, vol. 41, no. 2, Jun. 2022, doi: 10.6017/ital.v41i2.14683.
 - [131] S. Kohli, “Building User Trust in Conversational Ai: The Role of Explainable Ai in Chatbot Transparency,” Sep. 2024, doi: 10.5281/ZENODO.13833413.
 - [132] V. B. Nguyen, J. Schlötterer, and C. Seifert, “From Black Boxes to Conversations: Incorporating XAI in a Conversational Agent,” in *Explainable Artificial Intelligence*, vol. 1903, L. Longo, Ed., in Communications in Computer and Information Science, vol. 1903, Cham: Springer Nature Switzerland, 2023, pp. 71–96. doi: 10.1007/978-3-031-44070-0_4.
 - [133] D. Mindlin *et al.*, “Beyond one-shot explanations: a systematic literature review of dialogue-based xAI approaches,” *Artif Intell Rev*, vol. 58, no. 3, p. 81, Jan. 2025, doi: 10.1007/s10462-024-11007-7.
 - [134] B. T. Atmaja and A. Sasou, “Sentiment Analysis and Emotion Recognition from Speech Using Universal Speech Representations,” *Sensors*, vol. 22, no. 17, p. 6369, Aug. 2022, doi: 10.3390/s22176369.
 - [135] D. Resende Faria, A. I. Weinberg, and P. P. Ayrosa, “Multimodal Affective Communication Analysis: Fusing Speech Emotion and Text Sentiment Using Machine Learning,” *Applied Sciences*, vol. 14, no. 15, p. 6631, Jul. 2024, doi: 10.3390/app14156631.
 - [136] N. Elsayed, Z. ElSayed, N. Asadizanjani, M. Ozer, A. Abdelgawad, and M. Bayoumi, “Speech Emotion Recognition using Supervised Deep Recurrent System for Mental Health Monitoring,” in *2022 IEEE 8th World Forum on Internet of Things (WF-IoT)*, Yokohama, Japan: IEEE, Oct. 2022, pp. 1–6. doi: 10.1109/WF-IoT54382.2022.10152117.
 - [137] S. Mangalmurti, O. Saxena, and T. Singh, “Speech Emotion Recognition using CNN-TRANSFORMER Architecture,” in *2024 IEEE International Conference on Interdisciplinary Approaches in Technology and Management for Social Innovation (IATMSI)*, Gwalior, India: IEEE, Mar. 2024, pp. 1–6. doi: 10.1109/IATMSI60426.2024.10503276.
 - [138] Anjali Verma and Bappaditya Jnna, “Sentiment Analysis for Education and Behavior Analysis Based on Their Sentiment State,” *Int J Sci Res Sci & Technol*, vol. 12, no. 1, pp. 585–587, Feb. 2025, doi: 10.32628/IJSRST2411480.
 - [139] L. Kerkeni, Y. Serrestou, M. Mbarki, K. Raoof, and M. A. Mahjoub, “A review on speech emotion recognition: Case of pedagogical interaction in classroom,” in *2017 International Conference on Advanced Technologies for Signal and Image Processing (ATSIP)*, Fez, Morocco: IEEE, May 2017, pp. 1–7. doi: 10.1109/ATSIP.2017.8075575.
 - [140] M. Kastrati and M. Biba, “Natural language processing for Albanian: a state-of-the-art survey,” *IJECE*, vol. 12, no. 6, p. 6432, Dec. 2022, doi: 10.11591/ijece.v12i6.pp6432-6439.
 - [141] F. Kadriu, D. Murtezaj, F. Gashi, L. Ahmedi, A. Kurti, and Z. Kastrati, “Human-annotated dataset for social media sentiment analysis for Albanian language,” *Data in Brief*, vol. 43, p. 108436, Aug. 2022, doi: 10.1016/j.dib.2022.108436.
 - [142] K. Binjaku, J. Janku, and E. K. Mece, “Identifying Low-Resource Languages in Speech Recordings through Deep Learning,” in *2022 International Conference on Software, Telecommunications and Computer Networks (SoftCOM)*, Split, Croatia: IEEE, Sep. 2022, pp. 1–6. doi: 10.23919/SoftCOM55329.2022.9911376.
 - [143] D. Shehu, T. Luarasi, and P. Kyratsis, “Designing a natural language processing-driven communication system for urban planning: A case study,” *JCAU*, vol. 0, no. 0, p. 8417, Apr. 2025, doi: 10.36922/jcau.8417.
 - [144] E. Çano, “AlbNER: A Corpus for Named Entity Recognition in Albanian,” 2023, *arXiv*. doi: 10.48550/ARXIV.2309.08741.
 - [145] E. Çano, “AlbMoRe: A Corpus of Movie Reviews for Sentiment Analysis in Albanian,” 2023, *arXiv*. doi: 10.48550/ARXIV.2306.08526.

- [146] E. Çano and D. Lamaj, “AlbNews: A Corpus of Headlines for Topic Modeling in Albanian,” 2024, *arXiv*. doi: 10.48550/ARXIV.2402.04028.
- [147] Y. Kaloga and I. Kodrasi, “Towards interpretable emotion recognition: Identifying key features with machine learning,” 2025, *arXiv*. doi: 10.48550/ARXIV.2508.04230.
- [148] K. Venkataramanan and H. R. Rajamohan, “Emotion Recognition from Speech,” 2019, *arXiv*. doi: 10.48550/ARXIV.1912.10458.
- [149] Q. Li *et al.*, “Frame-Level Emotional State Alignment Method for Speech Emotion Recognition,” in *ICASSP 2024 - 2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, Seoul, Korea, Republic of: IEEE, Apr. 2024, pp. 11486–11490. doi: 10.1109/ICASSP48485.2024.10446812.
- [150] D. Tang, P. Kuppens, L. Geurts, and T. Van Waterschoot, “End-to-end speech emotion recognition using a novel context-stacking dilated convolution neural network,” *J AUDIO SPEECH MUSIC PROC.*, vol. 2021, no. 1, p. 18, Dec. 2021, doi: 10.1186/s13636-021-00208-5.
- [151] S. Madanian *et al.*, “Speech emotion recognition using machine learning — A systematic review,” *Intelligent Systems with Applications*, vol. 20, p. 200266, Nov. 2023, doi: 10.1016/j.iswa.2023.200266.
- [152] S.-G. Leem, D. Fulford, J.-P. Onnela, D. Gard, and C. Busso, “Selective Acoustic Feature Enhancement for Speech Emotion Recognition With Noisy Speech,” *IEEE/ACM Trans. Audio Speech Lang. Process.*, vol. 32, pp. 917–929, 2024, doi: 10.1109/TASLP.2023.3340603.
- [153] U. Bhattacharjee, S. Gogoi, and R. Sharma, “A statistical analysis on the impact of noise on MFCC features for speech recognition,” in *2016 International Conference on Recent Advances and Innovations in Engineering (ICRAIE)*, Jaipur, India: IEEE, Dec. 2016, pp. 1–5. doi: 10.1109/ICRAIE.2016.7939548.
- [154] F. Javanmardi, S. R. Kadiri, and P. Alku, “Pre-trained models for detection and severity level classification of dysarthria from speech,” *Speech Communication*, vol. 158, p. 103047, Mar. 2024, doi: 10.1016/j.specom.2024.103047.
- [155] P. Kumar, V. N. Sukhadia, and S. Umesh, “Investigation of Robustness of Hubert Features from Different Layers to Domain, Accent and Language Variations,” in *ICASSP 2022 - 2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, Singapore, Singapore: IEEE, May 2022, pp. 6887–6891. doi: 10.1109/ICASSP43922.2022.9746250.
- [156] W.-N. Hsu, B. Bolte, Y.-H. H. Tsai, K. Lakhotia, R. Salakhutdinov, and A. Mohamed, “HuBERT: Self-Supervised Speech Representation Learning by Masked Prediction of Hidden Units,” *IEEE/ACM Trans. Audio Speech Lang. Process.*, vol. 29, pp. 3451–3460, 2021, doi: 10.1109/TASLP.2021.3122291.
- [157] A. Chakhtouna, S. Sekkate, and A. Adib, “Unveiling embedded features in Wav2vec2 and HuBERT models for Speech Emotion Recognition,” *Procedia Computer Science*, vol. 232, pp. 2560–2569, 2024, doi: 10.1016/j.procs.2024.02.074.
- [158] H. Ji, T. Patel, and O. Scharenborg, “Predicting within and across language phoneme recognition performance of self-supervised learning speech pre-trained models,” 2022, *arXiv*. doi: 10.48550/ARXIV.2206.12489.
- [159] P. Tzirakis, G. Trigeorgis, M. A. Nicolaou, B. Schuller, and S. Zafeiriou, “End-to-End Multimodal Emotion Recognition using Deep Neural Networks,” 2017, doi: 10.48550/ARXIV.1704.08619.
- [160] M. El Ayadi, M. S. Kamel, and F. Karray, “Survey on speech emotion recognition: Features, classification schemes, and databases,” *Pattern Recognition*, vol. 44, no. 3, pp. 572–587, Mar. 2011, doi: 10.1016/j.patcog.2010.09.020.
- [161] A. Mohamed *et al.*, “Self-Supervised Speech Representation Learning: A Review,” *IEEE J. Sel. Top. Signal Process.*, vol. 16, no. 6, pp. 1179–1210, Oct. 2022, doi:

- 10.1109/JSTSP.2022.3207050.
- [162] S. Zhang, S. Zhang, T. Huang, and W. Gao, "Speech Emotion Recognition Using Deep Convolutional Neural Network and Discriminant Temporal Pyramid Matching," *IEEE Trans. Multimedia*, vol. 20, no. 6, pp. 1576–1590, Jun. 2018, doi: 10.1109/TMM.2017.2766843.
 - [163] N. Senthilkumar, S. Karpakam, M. Gayathri Devi, R. Balakumaresan, and P. Dhilipkumar, "Speech emotion recognition based on Bi-directional LSTM architecture and deep belief networks," *Materials Today: Proceedings*, vol. 57, pp. 2180–2184, 2022, doi: 10.1016/j.matpr.2021.12.246.
 - [164] A. Graves and J. Schmidhuber, "Framewise phoneme classification with bidirectional LSTM and other neural network architectures," *Neural Networks*, vol. 18, no. 5–6, pp. 602–610, Jul. 2005, doi: 10.1016/j.neunet.2005.06.042.
 - [165] D. Bahdanau, K. Cho, and Y. Bengio, "Neural Machine Translation by Jointly Learning to Align and Translate," 2014, *arXiv*. doi: 10.48550/ARXIV.1409.0473.
 - [166] S. Mirsamadi, E. Barsoum, and C. Zhang, "Automatic speech emotion recognition using recurrent neural networks with local attention," in *2017 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, New Orleans, LA: IEEE, Mar. 2017, pp. 2227–2231. doi: 10.1109/ICASSP.2017.7952552.
 - [167] RajratnaKharat*1, V. Bavane, S. Jadhao3, and R. Marode, "DIGITAL TWIN: MANUFACTURING EXCELLENCE THROUGH VIRTUAL FACTORY REPLICATION," Nov. 2018, doi: 10.5281/ZENODO.1493930.
 - [168] E. Glaessgen and D. Stargel, "The Digital Twin Paradigm for Future NASA and U.S. Air Force Vehicles," in *53rd AIAA/ASME/ASCE/AHS/ASC Structures, Structural Dynamics and Materials Conference & 20th AIAA/ASME/AHS Adaptive Structures Conference & 14th AIAA*, Honolulu, Hawaii: American Institute of Aeronautics and Astronautics, Apr. 2012. doi: 10.2514/6.2012-1818.
 - [169] D. Jones, C. Snider, A. Nassehi, J. Yon, and B. Hicks, "Characterising the Digital Twin: A systematic literature review," *CIRP Journal of Manufacturing Science and Technology*, vol. 29, pp. 36–52, May 2020, doi: 10.1016/j.cirpj.2020.02.002.
 - [170] F. Tao, M. Zhang, and A. Y. C. Nee, *Digital twin driven smart manufacturing*. London [England] ; San Diego, CA: Academic Press, 2019.
 - [171] A. Fuller, Z. Fan, C. Day, and C. Barlow, "Digital Twin: Enabling Technologies, Challenges and Open Research," *IEEE Access*, vol. 8, pp. 108952–108971, 2020, doi: 10.1109/ACCESS.2020.2998358.
 - [172] ASHRAE, "ASHRAE Technical FAQ 92: Recommended Indoor Temperature and Humidity Levels." Accessed: Apr. 13, 2026. [Online]. Available: https://www.ashrae.org/File%20Library/Technical%20Resources/Technical%20FAQs/TC-02.01-FAQ-92.pdf?utm_source=chatgpt.com
 - [173] CEN (European Committee for Standardization), *EN 16798-1:2019 Energy performance of buildings – Indoor environmental input parameters*, EN 16798-1:2019, 2019. Accessed: Apr. 13, 2026. [Online]. Available: <https://standards.iteh.ai/catalog/standards/cen/b4f68755-2204-4796-854a-6643dfcfe89/en-16798-1>
 - [174] E. R. Jones, J. G. Cedeño Laurent, A. S. Young, B. A. Coull, J. D. Spengler, and J. G. Allen, "Indoor humidity levels and associations with reported symptoms in office buildings," *Indoor Air*, vol. 32, no. 1, p. e12961, Jan. 2022, doi: 10.1111/ina.12961.

Lista e publikimeve

Artikujt shkencor

1. **Lili I.**, Kosta A., Xhina E., Mele R. (2025) “Integrating explainable AI (XAI) into Chatbots for automated decision-making in hospitality”, *Geojournal of Tourism and Geosites*, **Scopus**. <https://doi.org/10.30892/gtg.65209-1714> Year XIX, vol. 65, no. 2, 2026, p.719-729. E-ISSN 2065-1198. Botim në revistë shkencore ndërkombëtar.
2. Kosta A., **Lili I.**, Xhina E. (2025).”Analyzing Administrative Process Sequences Using PM4Py: A Case Study in an Albanian Municipality”, <https://doi.org/10.37745/ejsit.2013/vol13n51119137> *European Journal of Computer Science and Information Technology*, 13(51),119-137, 2025
3. **Lili I.**, and Kosta A. (2024) “Applying an Ordinary Least Squares (OLS) Regression Model on Processed Air Quality and Environment Data”, <https://doi.org/10.37745/bjes.2013/vol12n24958> *British Journal of Environmental Sciences* 12(2),49-58
4. Kosta A., **Lili I.**, and Xhina E. (2024) “Comparison and Evaluation of Air Quality Monitoring Methods Using IoT Devices”, <https://doi.org/10.37745/bjes.2013/vol12n3111> *British Journal of Environmental Sciences*, 12 (3), 1-11
5. Kosta A., **Lili I.**, Ceni A., Ktona A., Xhina E., Kika A., Loka E. (2026) A Review of AI Techniques for Emotion Recognition in Communication with Applications in Child Safety and Education, <https://doi.org/10.37745/ijeld.2013/vol14n2110> *International Journal of Education, Learning and Development*, 14 (2), 1-10

Konferenca

1. **Lili. I.**, Kosta. A., Xhina E. “The use of smart devices (IOT) to monitor the air quality: a case study at the Faculty of Natural Sciences” (2023) Conference: Recent Trends and Applications in Computer Science (RTA-CSIT): Faculty of Natural Sciences, Department of Informatics, Albania, ISSN 1613-0073
2. **Lili I.**, Ulqinaku M., Qosja F., Thomo A., Ktona A., Ceni A., Kosta A. (2025) “Sentiment Analysis from Speech in Albanian Language: A Deep Learning Approach for Low-Resource Scenarios”. *READ International Scientific Conference October 15-16, 2025: Tirana, Albania.. Publisher: Albanian-American Development Foundation with University of Tirana and University of Richmond*. Book of Abstract First Edition ISBN 9789928478849
3. **Lili I.**, Kosta A., Prifti G. (2025) “Labor Market Analysis of the IT Sector and Freelancing Using Web Scraping: A use case in Albania”. *8th International Conference on Applied Engineering and Natural Sciences ICAENS 2025 August 22-23: Konya, Turkey*. Proceeding book ISBN: 978-625-5794-24-6
4. **Lili I.**, Xhina E., and Kosta A. (2024) “Artificial Intelligence's role in recognizing and resolving student difficulties through personalized learning”. 1117-1123 *3rd International Conference on Engineering, Natural and Social Sciences May 16-17, 2024: Konya, Turkey. Proceeding book of 3rd international conference on engineering, natural and social sciences ICENSOS 2024*. Proceeding book ISBN: 978-625-6314-07-8

5. **Lili I.**, Xhina E., Kosta A., Ceni A., Ulqinaku M. (2024) “Comparing VADER and BERT for short-text Sentiment Analysis: Challenges and Observed Variations”. *3rd International Conference on Contemporary Academic Research ICCAR 2024 November 10-11, 2024: Konya, Turkey*. Proceeding book ISBN: 978-625-6314-68-9