

## ELEKTRON POCHTA TIZIMLARIDA INTELLEKTUAL XAVFSIZLIK VA QAROR QABUL QILISH

**Raximov Bobomurod Ramazonovich**

*TATU Mustaqil izlanuvchisi*

[Tel: +998997601554](tel:+998997601554)

### KIRISH

Zamonaviy axborot texnologiyalari davrida elektron pochta tizimlari foydalanuvchilar va korporativ muassasalar uchun asosiy kommunikatsiya vositasiga aylangan. Dunyo bo'ylab milliardlab foydalanuvchilar kundalik ishlarida elektron xatlardan foydalanadi, shu bilan birga ularning samaradorligi va xavfsizligi tizimga kiruvchi va chiquvchi ma'lumotlarning sifati va ishonchliligiga bog'liq. Biroq, elektron pochta tizimlarida kiberxavfsizlik tahdidlari, xususan, spam, phishing va zararli dasturlar tarqalishi tizimning ish faoliyatiga salbiy ta'sir ko'rsatadi, foydalanuvchi tajribasini pasaytiradi va moliyaviy hamda ma'lumot xavfsizligi sohasida jiddiy muammolarni yuzaga keltiradi. Shu sababli, elektron pochta tizimlarida intellektual xavfsizlik va qaror qabul qilish mexanizmlarini yaratish bugungi kunda dolzarb ilmiy va amaliy masala hisoblanadi.

An'anaviy qoidalarga asoslangan filtrlar ilgari spam va zararli xabarlarni aniqlashda foydali bo'lgan bo'lsa-da, ular murakkab va doimiy o'zgaruvchan xabar strukturalarini samarali qayta ishlay olmaydi. Shu bilan birga, so'nggi yillarda sun'iy intellekt (SI) va mashinani o'rganish (ML) usullaridan foydalangan holda elektron pochta xabarlarini avtomatik tasniflash, spam va phishing xabarlarini aniqlash, shuningdek, foydalanuvchi xatti-harakatlarini monitoring qilish imkoniyati sezilarli darajada oshdi [2]. Bu usullar xat matnidagi lingvistik va statistik belgilardan foydalangan holda, xabarlarni spam yoki no-spam sifatida avtomatik ajratish, tizimning o'z-o'zini moslashtirish qobiliyatini oshirish va real vaqt rejimida xavfsizlikni ta'minlash imkonini beradi.

Bugungi kunda mashinani o'rganishning turli algoritmlari, jumladan Naive Bayes, Support Vector Machines (SVM), Random Forest, sun'iy neyron tarmoqlar va transformer modellari, elektron pochta tizimlarida intellektual xavfsizlikni ta'minlashda keng qo'llaniladi. Har bir algoritmnining afzalliklari va kamchiliklari mavjud bo'lib, ularni samaradorlik, aniqlik, hisoblash xarajati va real vaqt rejimida ishlash nuqtai nazaridan solishtirish muhim ahamiyatga ega. Tadqiqotlar shuni ko'rsatadiki, chuqur o'rganish modellari va ansambl yondashuvlari murakkab xabarlarni aniqlashda yuqori aniqlik va F1 skor ko'rsatadi, klassik algoritmlar esa tezkor ishlash va kam resurs talab qilishi bilan ajralib turadi [1][2].

Tadqiqotning maqsadi; elektron pochta tizimlarida intellektual xavfsizlikni ta'minlash, spam va zararli xabarlarni aniqlash, shuningdek, foydalanuvchi va tizim xavfsizligini oshiradigan qaror qabul qilish mexanizmlarini ilmiy jihatdan tadqiq etishdan iborat. Tadqiqot natijalari amaliy jihatdan tizim samaradorligini oshirish, foydalanuvchi tajribasini yaxshilash va kiberxavfsizlikni ta'minlashga xizmat qilishi kutilmoqda.

### TADQIQOT METODOLOGIYASI

**Mazkur tadqiqotning maqsadi** — elektron pochta tizimlarida spam, phishing va zararli xabarlarini aniqlash hamda intellektual xavfsizlikni oshirish uchun sun'iy intellekt va mashinani o'rganish metodlarini tadqiq etishdir. Tadqiqot quyidagi bosqichlarda amalga oshiriladi:

Tadqiqotda ishlatiladigan asosiy ma'lumotlar elektron pochta xabarlaridan iborat bo'lib, ular spam va no spam (legitim) xabarlar sifatida belgilanadi. Ma'lumotlar to'plami quyidagi manbalardan olinadi: ochiq elektron pochta datasetlari, shuningdek, tadqiqot muallifi tomonidan to'plangan korporativ xabarlar. Ma'lumotlar tozalash bosqichida quyidagi jarayonlar amalga oshiriladi:

- duplikat xabarlarini olib tashlash;
- HTML teglar va keraksiz belgilarni tozalash;
- xabar matnlarini tokenizatsiya qilish va standartlashtirish;
- stop so'zlarni chiqarib tashlash va stemming lemmatizatsiya jarayonlarini qo'llash

#### **TAHLIL VA NATIJALAR**

Elektron pochta tizimlarida spam va zararli xabarlarini aniqlashda sun'iy intellekt (SI) va mashinani o'rganish (ML) yondashuvlari bugungi kunda eng samarali usullar sifatida qaraladi. Bu yondashuvlar klassik algoritmlardan tortib chuqur o'rganish modellarigacha bo'lgan keng qamrovli texnikalarni o'z ichiga oladi. Ularning samaradorligini baholashda odatda aniqlik (Accuracy), precision, recall va F1 skor kabi metrikalar qo'llaniladi, chunki bu ko'rsatkichlar modelning nafaqat to'g'ri klassifikatsiya qilish, balki noto'g'ri ijobiy va salbiy qarorlar ehtimolini ham o'lchaydi [3].

Klassik mashinani o'rganish algoritmlaridan bo'lgan Naive Bayes (NB), Support Vector Machines (SVM) va Random Forest (RF) algoritmlari ko'plab ilmiy tadqiqotlarda sinovdan o'tkazilgan. Masalan, bir tadqiqotda Naive Bayes 94,6 % aniqlik bilan ishlagan, SVM esa 95,2 % aniqlik va 94,8 % F1 skor bilan biroz yuqoriroq natija ko'rsatgan. Random Forest esa eng yaxshi natijani beradi, 96,8 % aniqlik va 96,2 % F1 skor bilan elektron pochta spamini aniqlashda yuqori samaradorlikni namoyon etgan [4].

Boshqa tadqiqotlarda SVM algoritmi juda yuqori natijalar ko'rsatgani qayd etilgan. Masalan, iKNiTO jurnalida chop etilgan maqolada SVM modeli 99,0 % aniqlik va 0,97 F1 skor bilan eng yuqori ko'rsatkichlardan birini ko'rsatgan, bu esa uning turli xil matnli xabarlarini umumlashtirish qobiliyati yuqori ekanligini tasdiqlaydi. Shu tadqiqotda Logistic Regression ham 98,4 % aniqlik bilan raqobatbardosh natija bergan, logistika regressiyasi modelining interpretatsiya imkoniyati tufayli xabar tuzilishining asosiy xususiyatlarini tahlil qilish qulayligi ta'kidlangan [3].

Transformer asosidagi chuqur o'rganish modellari, jumladan BERT, DistilBERT, RoBERTa va XLNet, kontekstni chuqur tahlil qilish orqali matnni samarali klassifikatsiya qila olganligi uchun muhim ahamiyat kasb etadi. Transformer modellarining 98,8 % aniqlik va 0,97 F1 skor bilan ishlashi, matnning semantic ma'nosini chuqur o'rganishning yuqori samaradorligini ko'rsatadi [5]. LSTM va boshqa RNN tarmoqlar ham spamni aniqlashda yuqori natijaga erishgan bo'lib, bu modellar kontekstni vaqt bo'yicha o'rganishda kuchli deb topilgan [3].

Bundan tashqari, turli tadqiqotlar ansambl texnikalari va gibrid modellarni ham o‘rganadi. Masalan, stacking yondashuvi yordamida bir nechta klassifikatorlarni birlashtirib ishlatish orqali aniqlik 98,8 % va F1 skor 98,9 % ga yetganligi ma’lum qilingan. Bu kabi ansambl usullar klassik modellar va yangi yondashuvlarning kuchli tomonlarini birlashtirish orqali yuqori natijalarga erishadi, shu bilan birga spamni aniqlashda barqarorlikni oshiradi [5].

Tahlil shuni ko‘rsatadiki, klassik ML algoritmlari va chuqur o‘rganish modellarining samaradorligi bir biridan farq qiladi, va ularning tanlovi ma’lumotlar to‘plami, hisoblash resurslari va real vaqtli tizimlar talablariga bog‘liq. Misol uchun, SVM kabi klassik modellarning yuqori aniqlik bilan ishlashi ularning nisbatan kichik va o‘rtacha hajmdagi datasetlar uchun juda mos ekanligini isbotlaydi, chunki SVM nozik chegaralarni topish orqali spam va no spam xabarlarini aniq ajratadi. Biroq, transformerlar kabi chuqur o‘rganish modellarining kontekstual bog‘liqliklarni chuqur tahlil qilishi murakkab semantik holatlarni aniqlashda ustunlik beradi [3][5].

Muvozanatni ta’minlash uchun tahlilda false positives va false negatives ning kamaytirilishi ham muhimdir. Spamni aniqlashda noto‘g‘ri ijobiy natija (legitim xabar spam deb belgilanishi) foydalanuvchi uchun katta noqulaylik tug‘diradi, shuning uchun F1 skor kabi metrikalar modelning balansli ishlashini baholash uchun muhim. Tadqiqotlarda transformer modellarining balansi yuqoriroq bo‘lishi, ularning nafaqat aniqlik, balki klassifikatsiya xavfsizligini ham oshirishi qayd etilgan [5].

Natijalar shuni ko‘rsatadiki, elektron pochta tizimlarida intellektual xavfsizlikni ta’minlashda SI asosida spam deteksiyasi tizimlari yuqori aniqlikka erishadi va bu tizimlar foydalanuvchi tajribasini yaxshilaydi hamda xavfsizlik tahdidlarini kamaytiradi [4][5]. Kelajakdagi tadqiqotlar bu modellarni real vaqt rejimida samarali ishlash, resurslarni kamaytirish va adversarial hujumlarga chidamliligini oshirish yo‘nalishlariga qaratilgan bo‘lishi mumkin [3][5].

## **XULOSA**

Tahlil natijalari shuni ko‘rsatadiki, ansambl va gibrid yondashuvlar klassik va chuqur o‘rganish modellarining afzalliklarini birlashtirib, aniqlik va F1 skor ko‘rsatkichlarini maksimal darajada oshiradi. Shu bilan birga, model tanlashda false positive va false negative holatlarining minimallashtirilishi, tizimning barqaror ishlashi va foydalanuvchi tajribasi ham muhim hisoblanadi.

Natijalar amaliy jihatdan elektron pochta tizimlarida intellektual xavfsizlikni ta’minlash, foydalanuvchilarni kiberxavfsizlik tahdidlaridan himoya qilish va tizim samaradorligini oshirishga xizmat qiladi. Kelajakdagi tadqiqotlar SI va mashinani o‘rganish modellarini real vaqt rejimida ishlash, resurslarni optimallashtirish va adversarial hujumlarga chidamlilikni oshirishga qaratilgan bo‘lishi lozim. Shu tarzda, elektron pochta tizimlarida qaror qabul qilishni avtomatlashtirish va xavfsizlikni intellektual darajada ta’minlash ilmiy va amaliy jihatdan dolzarb muammo sifatida qoladi.

**FOYDALANILGAN ADABIYOTLAR:**

1. Kumar S., Garg P. Spam Email Detection Using Machine Learning Techniques // International Journal of Computer Applications. – 2021. – Vol. 175, No. 30. – P. 25–33.
2. Ahmed M., Mahmood A. N., Hu J. A Survey of Network Anomaly Detection Techniques // Journal of Network and Computer Applications. – 2016. – Vol. 60. – P. 19–31.
3. IJICIS. Spam Detection in Emails Using Machine Learning Techniques: A Review. – 2023. – URL: [https://ijicis.journals.ekb.eg/article\\_406689.html](https://ijicis.journals.ekb.eg/article_406689.html)
4. Multiresearch Journal. Comparative Study of Machine Learning Algorithms for Email Spam Detection. – 2023. – URL: [https://www.multiresearchjournal.com/arclist/list-2025.5.4/id-4718?utm\\_source=.com](https://www.multiresearchjournal.com/arclist/list-2025.5.4/id-4718?utm_source=.com)
5. SpringerLink. Ensemble Methods for Spam Detection in Emails. – 2023. – URL: [https://link.springer.com/article/10.1007/s10207-023-00756-1?utm\\_source=.com](https://link.springer.com/article/10.1007/s10207-023-00756-1?utm_source=.com)

Global  
Science  
Publication