

Гульвира Бауржановна Бекешова – Л.Н. Гумилев атындағы Еуразия ұлттық университетінің аға оқытушысы, Қазақстан; e-mail: gulvirabauyrzhanovna@gmail.com. ORCID: <https://orcid.org/0000-0002-1635-4693>.

Information about the authors

Turaj Mehmanogly Mehdiyev – 2st year master's degree; specialty of information security; Eurasian National University named after L.N. Gumilyov; The Republic of Kazakhstan; e-mail: mehdiyev.t@gmail.com. ORCID: <https://orcid.org/0009-0004-6771-1584>.

Aigul Kairulayevna Shaykhanova – professor of the department of Information Security; Eurasian National University named after L.N. Gumilyov; Republic of Kazakhstan; e-mail: shaikhanova_ak@enu.kz. ORCID: <https://orcid.org/0000-0001-6006-4813>.

Gulvira Baurzhanovna Bekeshova – Senior Lecturer, L.N.Gumilev Eurasian National University, Kazakhstan; e-mail: gulvirabauyrzhanovna@gmail.com. ORCID: <https://orcid.org/0000-0002-1635-4693>.

Поступила в редакцию 17.10.2024

Поступила после доработки 21.10.2024

Принята к публикации 22.10.2024

[https://doi.org/10.53360/2788-7995-2024-4\(16\)-8](https://doi.org/10.53360/2788-7995-2024-4(16)-8)



МРНТИ: 81.93.29



А.Б. Какенова*, Б.К. Абдураимова, С.А. Сантеева
Л.Н. Гумилев атындағы Еуразия ұлттық университеті,
010008, Республика Казахстан, г. Астана, ул. Сатпаева, 2
*e-mail: akakenova1001@gmail.com

ЭЛЕКТРОНДЫҚ ПОШТА АРҚЫЛЫ ФИШИНГТІК ШАБУЫЛДАРДЫҢ АЛДЫН АЛУ ҮШІН ЖАСАНДЫ ИНТЕЛЛЕКТТІ ПАЙДАЛАНУ

Аңдатпа: Бұл мақалада электрондық пошта арқылы фишингтік шабуылдардың қазіргі проблемалары мен жойқын әсерлері қарастырылады. Ұйымдарды деректердің бұзылуынан және ықтимал апатты салдардан қорғау үшін алдын алу шараларының маңыздылығы атап өтілді. Фишингтік шабуылдардың негізгі түрлері және тәуекелдерді тиімді азайту үшін жасанды интеллект (AI) негізіндегі шешімдерді енгізу қажеттілігі сипатталған. Жұмыста жасанды интеллект пен машиналық оқытуды қолдана отырып, фишингтің алғашқы белгілерін тану әдістері қолданылды. Әзірлеу үшін Python және Google Colab қолданылды, бұл деректерді тиімді талдауға және модельдерді оқытуға мүмкіндік берді. Ерекше әдістерді әзірлеуге және заманауи бағдарламалық жасақтаманы пайдалануға ерекше назар аударылды. Зерттеу нәтижесінде фишингтік шабуылдарды танудағы AI құралдарының жоғары тиімділігін растайтын деректер алынды. AI технологиялары фишингті анықтаудың дәлдігін едәуір арттырды және жаңа киберқауіптерге бейімделуді қолдады. Талдау көрсеткендей, AI қолдану шабуылдарды уақтылы анықтауға ғана емес, сонымен қатар алдын алу стратегияларын жасауға мүмкіндік береді. Нәтижелердің практикалық маңыздылығы-әзірленген әдістерді қолданыстағы қауіпсіздік жүйелеріне біріктіру мүмкіндігі. Жұмыс ақпаратты тұрақты қорғауды құру үшін технологиялық жетістіктер мен ұйымдастырушылық тәжірибелерді біріктіретін стратегиялық тәсілді ұсынады.

Түйін сөздер: спам, фишинг, ақпараттық қауіпсіздік, машиналық оқыту, сүзу, әдістер, қорғау, қиындықтар.

Кіріспе

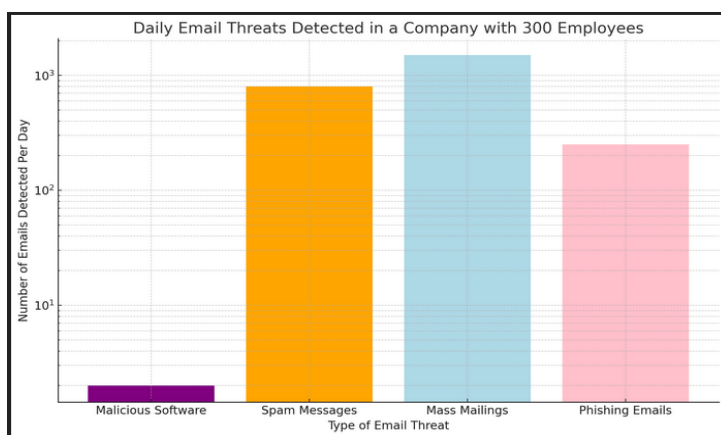
Электрондық пошта фишингі-бұл алаяқтық әдісі, онда жіберушілер өз хабарламаларын заңды сұраулар ретінде жасырады, олардың мақсаты алушылардың жеке мәліметтерін алу үшін алдау болып табылады. Жәбірленушілер қауіп-қатерден бейхабар болып, сақтандыру полистерін, банктік шоттар және әлеуметтік сақтандыру деректер сияқты құпия ақпаратты

ұсынады. Бұл шабуылдар бүкіл әлемдегі кәсіпорындарға таралатын және олардың талғампаздығын көрсететін ұйымдар үшін айтарлықтай қауіп төндіреді.

Электрондық пошта арқылы фишингтік науқандар олардың ықтимал жойқын салдарын ескере отырып, үлкен қауіп ретінде қарастырылады. Қаржылық шығындар айсбергтің шыңы ғана; терең шығындар маңызды ақпараттың ағып кетуін қамтиды, компанияны айтарлықтай заңды тәуекелдерге ұшыратады және клиенттердің сеніміне нұқсан келтіреді. Бұл корпоративтік поштаны фишингтік шабуылдардан қорғау үшін кешенді шараларды енгізу қажеттілігін көрсетеді [1].

Электрондық поштаны талдауға негізделген фишингтік шабуылдардың статистикасы мен ауқымы.

Жаңа өнімді шығаруға дайындық жұмыстары аясында Kaspersky Lab сарапшылары үш миллионнан астам электрондық поштаны талдады. Талдау нәтижесінде зиянды бағдарламалық жасақтаманың 800-ге жуық данасы, 60 мың спам-хабарлама, 110 мың жаппай пошта және 20 мың фишингтік хат табылды. Иллюстрация үшін біз күніне орта есеппен 40 мың хабарлама алатын 300 адамнан тұратын компаниядағы жағдайды ұсынылған. Бұл электрондық поштаны пайдаланудың шамамен 2,5 айында аталған үш миллион электрондық пошта жиналғанын білдіреді. Бір жұмыс күні ішінде, мұндай компания 800 спам-хабарламамен, шамамен 1500 жаппай жіберіліммен және 250-ден астам фишингтік хаттармен бетпе-бет келеді, зиянды бағдарлама анықталған 1-2 жағдайды айтпағанда, мұның бәрі пошта Office 365-ке енгізілген қорғаныс құралдарымен сүзілгеннен кейін пайда болады, олар кездейсоқ 3,9% қажетті хаттар жояды, оларды зиянды деп қателескен аналитикасын сурет 1-де көрсетілген.



Сурет 1 – Kaspersky Lab аналитикасы

Жасанды интеллектті оқыту: әдістері мен алгоритмдері

Жасанды интеллект (AI) тренингі-бұл машинаның деректер мен ережелер жиынтығын қолдана отырып, белгілі бір тапсырмаларды орындауды меңгеру процесі. Оқытудың әртүрлі әдістері мен алгоритмдері мәліметтерден ақпарат алуға және негізделген шешімдер қабылдауға мүмкіндік береді [2]. Бұл мақалада біз AI оқытудың негізгі әдістері мен алгоритмдерін қарастырылады, сонымен қатар оларды Python-да қолдануға мысалдар келтіріледі.

Мұғаліммен оқыту (Supervised Learning) әдісі модельді оқыту үшін орналастырылған деректерді (мысалы, жіктелген кескіндер) пайдалануды қамтиды. Алгоритмдердің мысалдарына сызықтық регрессия, логистикалық регрессия және шешім ағаштары жатады[2]. Бұл әдіс жіктеу, регрессия және болжау есептері үшін қолданылады. Python-дағы мысалы сурет 2 – сызықтық регрессия көрсетілген.

Мұғалімсіз оқыту әдісі тегтерсіз деректерді талдау үшін қолданылады (мысалы, жіктелмеген кескіндер). Алгоритмдердің мысалдарына кластерлеу, негізгі компоненттерді талдау және факторлық талдау жатады. Бұл әдіс деректер құрылымын анықтау, ауытқуларды анықтау және деректерді қысу үшін қолданылады.

Белсенді оқыту әдісі модельге қазіргі білімі негізінде оқу үшін деректерді өз бетінше таңдауға мүмкіндік береді. Алгоритмдердің мысалдарына ең Ақпараттық нүктелерді таңдау

және модельді қолдану әдістері жатады. Бұл әдіс деректерді жинауға және өңдеуге уақыт пен ресурстарды үнемдейді. Python-дағы мысал сурет 3 Modal модульдік кітапханасын қолдана отырып белсенді оқыту көрсетілген.

```
from sklearn.model_selection import train_test_split
from sklearn.linear_model import LinearRegression
from sklearn.metrics import mean_squared_error
import pandas as pd

# Uploading data
data = pd.read_csv('data.csv')
X = data[['feature1', 'feature2']]
y = data['target']

# Separation of data into training and test data
X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2, random_state=42)

# Creating and training a model
model = LinearRegression()
model.fit(X_train, y_train)

# Prediction and evaluation of the model
y_pred = model.predict(X_test)
print("Mean Squared Error:", mean_squared_error(y_test, y_pred))
```

Сурет 2 – Сызықтық регрессия функциясы

```
[ ] from sklearn.datasets import make_classification
from sklearn.model_selection import train_test_split
from sklearn.ensemble import RandomForestClassifier
from modal.models import ActiveLearner
from modal.uncertainty import uncertainty_sampling

# Creating synthetic data
X, y = make_classification(n_samples=1000, n_features=20, n_informative=2, n_redundant=2, random_state=42)

# Separation of data into training and test data
X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2, random_state=42)

# Initializing an active student
learner = ActiveLearner(
    estimator=RandomForestClassifier(),
    query_strategy=uncertainty_sampling,
    X_training=X_train, y_training=y_train
)

# Training the model on the initial data set
learner.fit(X_train, y_train)

# Sampling new training data
query_idx, query_instance = learner.query(X_test)
learner.teach(X_test[query_idx].reshape(1, -1), y_test[query_idx].reshape(1, ))

#Evaluation of the model
print("Model accuracy:", learner.score(X_test, y_test))
```

Сурет 3 – Modal модульдік кітапханасын қолдана отырып белсенді оқыту

AI оқыту деректерден білім алу үшін әртүрлі әдістер мен алгоритмдерді пайдаланады. Әдісті таңдау тапсырма түріне, деректер көлеміне және қол жетімді ресурстарға байланысты. Айта кету керек, AI-ді оқыту тек осы әдістермен ғана шектелмейді және үнемі жаңа әдістер мен алгоритмдер жасалуда[3].

Электрондық пошта мекенжайларын жіктеуге арналған машиналық оқыту алгоритмдері

Жасанды интеллектке негізделген электрондық поштаны сүзу жүйелері негізінен машиналық оқыту алгоритмдерінің екі түрін пайдаланады. Бұл мұғаліммен машиналық оқытудың жіктеуіштері және мұғалімсіз машиналық оқыту алгоритмдері. Бақыланатын машиналық оқыту классификаторларын маркетингтік электрондық пошталар, іскери байланыстардағы маңызды электрондық хаттар, қаржы институттарының шынайы хаттары, спам-хабарламалар және фишингтік хаттар сияқты әртүрлі санаттарға тән электрондық пошта сипаттамаларын автоматты түрде анықтауға үйретуге болады [4]. Қолмен жасалған ережелерге тәуелді дәстүрлі ережелерге негізделген электрондық пошта сүзгілерінен айырмашылығы, бақыланатын жіктеуіштер фишингтік хаттардың жаңа мысалдарын зерттей отырып, спам мен фишингтің жаңа тенденцияларына үнемі бейімделіп отырады.

Екінші жағынан, машиналық оқытудың бақыланбайтын алгоритмдерін заңды электрондық поштаның профилін жасау үшін пайдалануға болады. Бұл әдістер пайдаланушы мен бизнестің мінез-құлқы мен өзара әрекеттесуіне негізделген электрондық пошта

профильдерін құру арқылы жұмыс істейді. Олар негізінен ішкі қауіптерді анықтау үшін қолданылады.

Электрондық пошта мекенжайларын жіктеуге арналған машиналық оқыту алгоритмдері ақпараттық қауіпсіздікті нығайтуда аса маңызды роль атқарады. Жіктеуіштерді дамыту барысында бақыланатын машиналық оқыту модельдері үлгілерді тану және сараптау арқылы фишингтік және спам хаттарды анықтайды. Осындай модельдер адам қолымен белгіленген электрондық пошталардан үйренеді және қауіп төндіретін хаттардың жаңа үлгілерін ажырата алады. Бұлардың көмегімен, спамның жаңа нұсқаларына бейімделуге және қолданушыларға бағытталған шабуылдарды болдырмауға мүмкіндік туады[5].

Бақыланбайтын оқыту алгоритмдері пайдаланушының өзара әрекеттесуі мен электрондық поштаның мазмұнын талдау арқылы, ішкі қауіптерді анықтау үшін электрондық поштаның өздігінен дамитын профилін құрады. Мұндай профильдер қолданушылардың әдеттегі хат алмасу өрнектерінен ауытқуларды байқап, аномалияларды тез анықтауға көмектеседі.

Ақпараттық қауіпсіздікті күшейту стратегиясы ретінде, жасанды интеллект және машиналық оқыту алгоритмдерін қолдану үшін әртүрлі тәсілдер мен модельдерді біріктіру қажет. Нейрондық желілер, шешім ағаштары, логистикалық регрессия және тағы басқалар сияқты әдістерді қолдану арқылы, жүйе қауіпті хаттардың әртүрлі сипаттамаларын тануға үйренеді. Осы машиналық оқыту әдістерінің артықшылығы – олардың өзгермелі мәліметтерге бейімделуі, жаңа түрдегі шабуылдар мен спам әдістерін тез анықтауы және қолданушылардың әрекеттерінен үлгілерді шығара алуы. Осы қасиеттердің барлығы электрондық пошталардың сүзгілеу жүйесінің тиімділігін арттырады. Көптеген жағдайларда, бұл алгоритмдер белгілі бір мәліметтер жиынтығына тәуелсіз жұмыс істейді, бұл оларды әмбебап. Қазіргі заманғы электрондық пошта сүзгілерінің қабілеттілігі ақпараттық қауіпсіздігін арттыруда маңызды рөл атқарады. Бақыланатын машиналық оқыту алгоритмдері және бақыланбайтын әдістер ақпараттың ағып кетуі мен киберқауіпсіздіктің басқа да қатерлеріне қарсы күресуде негізгі құралға айналууда. Бұл жүйелердің арқасында, ұйымдар спаммен, вирустармен және басқа да зиянды бағдарламалармен күресуде белсенді түрде әрекет ете алады, өз қызметкерлерінің ақпараттық қауіпсіздік мәдениетін нығайтады [6].



Сурет 4 – Желі архитектурасы

Хабарламаларды бақылау және спамды сүзу үшін жасанды интеллект құру бірнеше қадамдарды қамтиды. Біріншіден, желі архитектурасына сүйеніп, база құрылады, бұл сурет 4 – Желі архитектурасы көрсетілген. Міне, жалпы процесс:

- Деректерді жинау: «Спам» және «спам емес» белгілері бар деректерді жиналады. Бұл деректер мәтіндік хабарларды, электрондық пошталарды және басқа хабар түрлерін қамтуы мүмкін.

- Деректерді өңдеу: қажет емес таңбаларды, HTML тегтерін және басқа қоқыстарды жою арқылы деректерді тазалайды. Мәтінді сандық форматқа түрлендіріледі, мысалы, TF-IDF

(term Frequency-Inverse Document frequency) әдістері немесе сөз векторлары (Word Embeddings).

- Деректерді бөлу: деректерді оқу және сынақ үлгілеріне бөлінеді (сурет 5 – Фишингтік шабуылдарды анықтау функциясы).

- Модель құру: логистикалық регрессия, шешім ағаштары, кездейсоқ ормандар немесе нейрондық желілер сияқты жіктеу үшін машиналық оқыту алгоритмін таңдалады. Модельді деректерді оқыту үлгісінде оқытылады.

- Модельді бағалау: сынақ үлгісіндегі модельді дәлдік, толықтық, F1 өлшемдері және басқаларын пайдаланып бағаланады.

- Мониторинг жүйесін әзірлеу: электрондық пошта қызметімен немесе мессенджермен біріктірілетін жүйені енгізіледі. Спам ретінде жіктелген хабарламалар спам қалтасына көшіріледі [7].

Мұндай жүйені нақты электрондық пошта қызметтерімен немесе мессенджерлермен біріктіру үшін сізге осы қызметтердің API интерфейстерін пайдалану қажет болады. Мысалы, Gmail үшін бұл Google API, Telegram үшін – Telegram Bot API және т.б. болуы мүмкін. Фишингтік шабуылдарды анықтау функциясын қосу үшін модель құру процесіне қосымша қадамдар қосу қажет. Біз модельді спамды ғана емес, фишингтік хабарламаларды да тануға үйретіледі, сурет 5 – фишингтік шабуылдарды анықтау функциясы көрсетілген. Мұны деректерге жаңа фишинг белгісін қосу арқылы жасауға болады.

```
{x}
# Usage example
message = "Вы выиграли миллион долларов!"
classification = classify_message(message)

if classification == 'spam':
    print("Это спам!")
elif classification == 'phishing':
    print("Это фишинг!")
else:
    print("Это не спам.")

# For integration with the system
def process_messages(messages):
    results = []
    for msg in messages:
        classification = classify_message(msg)
        if classification == 'spam':
            results.append((msg, 'Перемещено в папку спам'))
        elif classification == 'phishing':
            results.append((msg, 'Перемещено в папку фишинг'))
        else:
            results.append((msg, 'Оставлено во входящих'))
    return results

# Example of processing a list of messages
messages = [
    "Поздравляем! Вы выиграли бесплатный iPhone. Перейдите по ссылке...",
    "Ваша учетная запись заблокирована. Сбросьте пароль, перейдя по ссылке...",
    "Привет, как дела?"
]

results = process_messages(messages)
for msg, action in results:
    print(f'Сообщение: "{msg}" - Действие: {action}')
```

Сурет 5 – Фишингтік шабуылдарды анықтау функциясы

```
from sklearn.linear_model import LogisticRegression
from sklearn.metrics import classification_report, accuracy_score

# Модельді дайындау
model = LogisticRegression()
model.fit(X_train_counts, df_emails_train['label'])

# Модельді бағалау
X_test_counts = vectorizer.transform(df_emails_test['message'])
y_pred = model.predict(X_test_counts)

print(classification_report(df_emails_test['label'], y_pred))
print(f'Accuracy: {accuracy_score(df_emails_test['label'], y_pred)}')
```

Сурет 6 – Модельді оқыту және бағалау

Осы код арқылы біз деректер жиынтығын біріктіреді, оны оқыту және тестілеу үлгілеріне бөледі, бұл сурет 6 – Модельді оқыту және бағалау көрсетілген, сөздердің жиілік таралуын талдайды және логистикалық регрессияны қолдана отырып, модель дайындалады. Жиілік талдауы фишингтік хаттарда жиі қолданылатын кілт сөздерді анықтауға көмектеседі, бұл фишингті анықтау механизмдерін жақсартуға мүмкіндік береді [5].

Шешімдерді дайындау кезінде талданған нақты деректердің мысалдары.

Алдымен датасет, `pd.concat()` функциясы- python-дағы `pandas` кітапханасының бөлігі, ол `pandas` нысандарын берілген ось бойымен біріктіруге арналған. Мысалда, датасет `ham`

және датасет phish біріктірілді, бұл сурет 7 де көрсетілген. Бұл код электрондық пошта деректерінің жиынтығын оқыту және тестілеу үлгілеріне бөлуді жүзеге асырады, сонымен қатар электрондық пошта мазмұнындағы сөздердің жиіліктік таралуын талдайды. Алдымен train_test_split функциясын қолдана отырып бөлу жасалады, содан кейін тоқтату сөздерін, сондай-ақ "font" және "subject" сөздерін қоспағанда, оқу жиынтығындағы барлық сөздердің жиіліктік таралуы жасалады. Осыдан кейін фишингтік хаттар үшін сөздердің жиіліктік таралуы құрылады және ең көп кездесетін 20 сөздің графигі көрсетіледі [8].

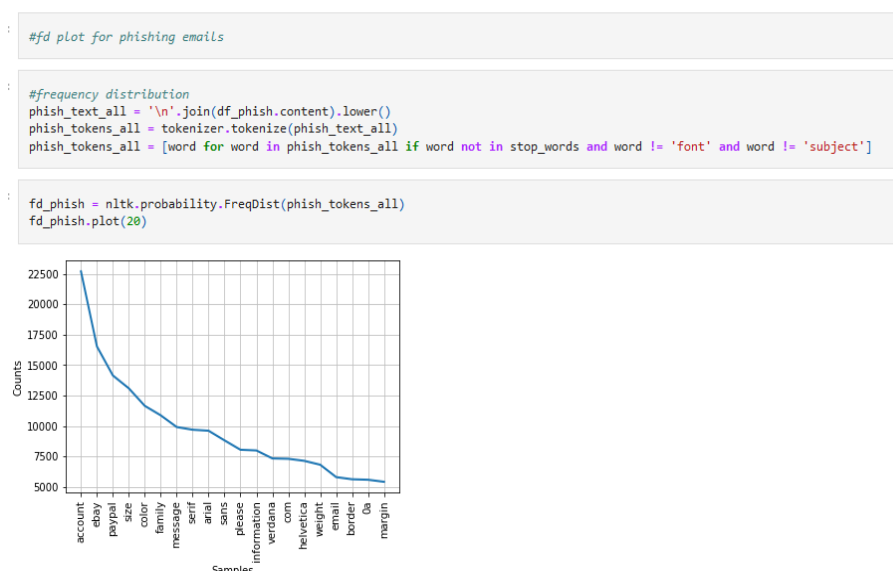
Код әрекеттерінің сипаттамасын қысқаша қайта құрайық: train_test_split функциясын қолдана отырып, df_emails электрондық поштасының деректер жиынтығы Оқу (df_email_train) және тест (df_emails_test) үлгілеріне бөлінеді, мұнда деректердің 20% тестілеуге бөлінеді. Әріптерде жиі қолданылатын сөздер туралы жалпы түсінік алу үшін тоқтату сөздері мен 'қаріп' және 'тақырып' сөздерін қоспағанда, Оқу үлгісі үшін сөздердің жиілік таралуы жасалады [9].

Код әрекеттерінің сипаттамасын қысқаша қайта құрайық: train_test_split функциясын қолдана отырып, df_emails электрондық поштасының деректер жиынтығы Оқу (df_email_train) және тест (df_emails_test) үлгілеріне бөлінеді, мұнда деректердің 20% тестілеуге бөлінеді. Әріптерде жиі қолданылатын сөздер туралы жалпы түсінік алу үшін тоқтату сөздері мен 'қаріп' және 'тақырып' сөздерін қоспағанда, Оқу үлгісі үшін сөздердің жиілік таралуы жасалады [9].

```
df_emails = pd.concat([df_ham, df_phish])
df_emails.sample(10)
```

	label		title	content
2526	ham	4010.2001-11-15.kitchen.ham.txt	Subject: org chart\ntina rode\ncassitant to da...	
738	phish	377_False.txt	Subject: New email address added to your accou...	
2045	phish	1528_False.txt	Subject: Account Review PayPal <\n-- hr\ndotte...	
710	ham	1678.2001-07-30.williams.ham.txt	Subject: party @ miss amy 's house fri\ntwe wo...	
3818	ham	0222.2001-03-16.kitchen.ham.txt	Subject: re : inland resources\ndick - great j...	
200	ham	4646.2002-01-12.williams.ham.txt	Subject: start date : 1 / 12 / 02 ; hourahead ...	
3030	ham	1836.2001-07-30.kitchen.ham.txt	Subject: beaver creek accommodations\ntbelow is...	
1174	ham	3982.2001-11-13.kitchen.ham.txt	Subject: today 's floor meeting\ntyou may get ...	
508	phish	2180_False.txt	Subject: About your account\nt You have added s...	
111	phish	862_False.txt	Subject: Notification from Billing Department ...	

Сурет 7 – Google Colab-та ашылған датасет



Сурет 8 – Датасетті санаттар бойынша кездесу жиіліктен бөлінген график

Фишингтік хаттар үшін (df_phish) тоқтату сөздері мен қалаусыз «қаріп» және «тақырып» сөздерін сүзгеннен кейін сөздердің жиіліктік таралуы жасалады (сурет 8. Датасетті санаттар бойынша кездесу жиіліктен бөлінген график). Жиілікті талдау нәтижелерін визуализациялау үшін фишингтік хаттарда жиі кездесетін 20 сөздің графигі жасалады. Бұл талдау фишингтік

хаттарда жиі қолданылатын кілт сөздерді анықтауға көмектеседі, бұл фишингті анықтау механизмдерін жақсартуға көмектеседі.

Қорытынды

Жасанды интеллект спамды, фишингті және басқа да зиянды шабуылдарды анықтауға және бұғаттауға арналған қуатты құралдарды ұсына отырып, заманауи поштаны сүзу жүйелерінде шешуші рөл атқарады. Машиналық оқытудың әртүрлі әдістеріне, сондай-ақ табиғи тілді өңдеу әдістеріне сүйене отырып, кіріс поштаны автоматты түрде талдауға және жіктеуге қабілетті тиімді алгоритмдер жасалады [10]. Поштаны сүзу үшін жасанды интеллектті пайдалану зиянды хабарларды анықтау және бұғаттау процесін жақсартуға, ұйымға сәтті шабуыл жасау мүмкіндігін азайтуға және электрондық поштаның өнімділігі мен тиімділігін арттыруға мүмкіндік береді. Алайда, технологияның дамуы қауіпсіздік шараларын айналып өту әдістерін жетілдіруге әкелетінін түсіну қажет. Сондықтан сүзу алгоритмдерін үнемі жетілдіріп отыру, шабуылдардың жаңа түрлерін қадағалау және киберқауіпсіздік саласындағы өзгеріп отыратын қауіптерге бейімделу маңызды. Жалпы, поштаны сүзу үшін жасанды интеллектті қолдану киберқауіптермен күресуде және пайдаланушы деректерінің құпиялылығын қорғауда тиімді құрал болып табылады.

Болашақ бағыттар мен қиындықтар

Спам классификациясындағы айтарлықтай жетістіктерге қарамастан, шешілмеген мәселелер мен жаңа қиындықтар әлі де бар. Болашақ зерттеулердің негізгі бағыттарының бірі- нақты уақыт режимінде жұмыс істей алатын жүйелерді дамыту. Қазіргі жағдайда жіктеудің жылдамдығы мен тиімділігі, әсіресе корпоративті пайдаланушылар мен ірі ұйымдар үшін өте маңызды болады. Сонымен қатар, жүйенің функцияларын динамикалық жаңартуға назар аудару керек, бұл жүйені толығымен қайта құруды қажет етпестен спамерлердің жаңа әдістеріне бейімделуге мүмкіндік береді.

Маңызды хаттар спамға түспеуі үшін жалған позитивтерді азайту да маңызды. Жіктеудің жоғары дәлдігі контекстік талдаудың терең интеграциясын және табиғи тілді өңдеудің заманауи технологияларын қолдануды қажет ететін жалған оң нәтижелердің төмен деңгейімен бірге жүруі керек. Спамды супервайзерден бақылаусыз оқытуға жіктеу әдістерінің эволюциясы икемді, дәл және тиімді жүйелерді құруға деген ұмтылысты көрсетеді. Супервайзерлік әдістер ең сенімді және дәл болып қалса да, супервайзер емес және жартылай бақылаушы әдістер тез өзгертін жағдайлар мен қауіптерге бейімделудің жаңа мүмкіндіктерін ұсынады. Осы саладағы болашақ зерттеулер нақты уақыттағы жүйелерді жобалауға, мүмкіндіктерді динамикалық жаңартуға және спамды сүзудің неғұрлым сенімді және тиімді жүйелерін құруға мүмкіндік беретін жалған позитивтердің деңгейін төмендетуге назар аударуы керек.

Әдебиеттер тізімі

1. Абрамов И.П. Использование искусственного интеллекта в системах фильтрации электронной почты. / И.П. Абрамов, А.Ю. Бородин // Информационные технологии и телекоммуникации. – 2021. – № 1(32). – С. 59-66.
2. Маленков М.Г. Защита от атак с использованием ИИ с помощью ИИ / М.Г. Маленков // Инновации. Наука. Образование. – 2021. – Том 2, № 44. – с. 49-53.
3. Исаков А.А. Искусственный интеллект и расследование киберпреступлений / А.А. Исаков // Вестник науки. – 2023. – № 5(62). – С. 597-603.
4. Трофимов А.И. Использование методов искусственного интеллекта для борьбы с фишингом и спамом в электронной почте / А.И. Трофимов // Компьютерные инструменты в образовании. – 2022. – № 26(1). – С. 49-55.
5. Миронов А.А. Применение методов машинного обучения и искусственного интеллекта для защиты от спама и фишинга / А.А. Миронов // Системы и инструменты информатики. – 2020. – № 30(2). – С. 57-69.
6. Sharma A. Email Spam Filtering Using Machine Learning Algorithms: A Review. / A. Sharma, R. // Gupta In Advances in Computer Vision & Image Processing. – 2021. – Vol. 2. – P. 137-145.
7. Пожогин А. Проблемы облачной почты: спам, фишинг и вредоносное ПО [электронный ресурс] <https://blog.kaspersky.kz/spam-phishing-malware/1804/>.
8. Sh.Drew, Phishing-Detection, [elektronnyj resurs] <https://github.com/shaferd/Phishing-Detection/>.

9. Кузнецов Д.А. Использование нейронных сетей для фильтрации фишинга / Д.А. Кузнецов // Журнал информационных систем. – 2023. – № 6(41). – С. 110-118.
10. Васильев Р.Н. Анализ эффективности алгоритмов классификации электронной почты / Р.Н. Васильев, И.К. Петрова // Технологии защиты информации. – 2023. – № 2(39). – С. 66-73.

References

1. Abramov I.P. Ispol'zovanie iskusstvennogo intellekta v sistemakh fil'tratsii ehlektronnoi pochty. / I.P. Abramov, A.YU. Borodin // Informatsionnye tekhnologii i telekommunikatsii. – 2021. – № 1(32). – S. 59-66. (In Russian).
2. Malenkov M.G. Zashchita ot atak s ispol'zovaniem II s pomoshch'yu II / M.G. Malenkov // Innovatsii. Nauka. Obrazovanie. – 2021. – Tom 2, № 44. – S. 49-53. (In Russian).
3. Isakov A.A. Iskusstvennyi intellekt i rassledovanie kiberprestuplenii / A.A. Isakov // Vestnik nauki. – 2023. – № 5(62). – S. 597-603. (In Russian).
4. Trofimov A.I. Ispol'zovanie metodov iskusstvennogo intellekta dlya bor'by s fishingom i spamom v ehlektronnoi pochte / A.I. Trofimov // Komp'yuternye instrumenty v obrazovanii. – 2022. – № 26(1). – S. 49-55. (In Russian).
5. Mironov A.A. Primenenie metodov mashinnogo obucheniya i iskusstvennogo intellekta dlya zashchity ot spama i fishinga / A.A. Mironov // Sistemy i instrumenty informatiki. – 2020. – № 30(2). – S. 57-69. (In Russian).
6. Sharma A. Email Spam Filtering Using Machine Learning Algorithms: A Review. / A. Sharma, R. // Gupta In Advances in Computer Vision & Image Processing. – 2021. – Vol. 2. – R. 137-145. (In English).
7. Pozhogin A. Problemy oblachnoi pochty: spam, fishing i vredonosnoe PO [ehlektronnyi resurs] <https://blog.kaspersky.kz/spam-phishing-malware/1804/>. (In Russian).
8. Sh.Drew, Phishing-Detection, [elektronnyj resurs] <https://github.com/shaferd/Phishing-Detection/>. (In English).
9. Kuznetsov D.A. Ispol'zovanie neironnykh setei dlya fil'tratsii fishinga / D.A. Kuznetsov // Zhurnal informatsionnykh sistem. – 2023. – № 6(41). – S. 110-118. (In Russian).
10. Vasil'ev R.N. Analiz ehffektivnosti algoritmov klassifikatsii ehlektronnoi pochty / R.N. Vasil'ev, I.K. Petrova // Tekhnologii zashchity informatsii. – 2023. – № 2(39). – S. 66-73. (In Russian).

А.Б. Какенова*, Б.К. Абдураимова, С.А. Сантеева
Л.Н.Гумилев атындағы Еуразия ұлттық университеті,
010008, Республика Казахстан, г. Астана, ул. Сатпаева, 2
*e-mail: akakenova1001@gmail.com

ИСПОЛЬЗОВАНИЕ ИИ ДЛЯ ПРЕДОТВРАЩЕНИЯ ФИШИНГОВЫХ АТАК ПО ЭЛЕКТРОННОЙ ПОЧТЕ

В этой статье рассматриваются текущие проблемы и разрушительные последствия фишинговых атак по электронной почте. Подчеркивается важность профилактических мер для защиты организаций от утечки данных и потенциальных катастрофических последствий. Описаны основные типы фишинговых атак и необходимость внедрения решений на основе искусственного интеллекта (ИИ) для эффективного снижения рисков. В работе использовались методы распознавания ранних признаков фишинга с использованием искусственного интеллекта и машинного обучения. Для разработки использовались Python и Google Colab, что позволило эффективно анализировать данные и обучать модели. Особое внимание было уделено разработке уникальных методов и использованию современного программного обеспечения. В результате исследования были получены данные, подтверждающие высокую эффективность средств ИИ в распознавании фишинговых атак. Технологии искусственного интеллекта значительно повысили точность обнаружения фишинга и поддержали адаптацию к новым киберугрозам. Анализ показывает, что использование ИИ позволяет не только своевременно выявлять атаки, но и разрабатывать стратегии профилактики. Практическая значимость результатов заключается в возможности интеграции разработанных методов в существующие системы безопасности. Работа предлагает стратегический подход, сочетающий технологические достижения и организационные практики для создания устойчивой защиты информации.

Ключевые слова: спам, фишинг, информационная безопасность, машинное обучение, фильтрация, методы, защита, проблемы.

A.B. Kakenova*, B.K. Abduraimova, S.A. Santeeva
L.N. Gumilyov Eurasian National University,
010008, Republic of Kazakhstan, Astana, street Satpayeva, 2
*e-mail: akakenova1001@gmail.com

USING AI TO PREVENT EMAIL PHISHING ATTACKS

This article examines the current problems and the devastating effects of phishing email attacks. The importance of preventive measures to protect organizations from data leakage and potential catastrophic consequences is emphasized. The main types of phishing attacks and the need to implement solutions based on artificial intelligence (AI) to effectively reduce risks are described. The work used methods for recognizing early signs of phishing using artificial intelligence and machine learning. Python and Google Colab were used for development, which made it possible to effectively analyze data and train models. Special attention was paid to the development of unique methods and the use of modern software. As a result of the study, data were obtained confirming the high effectiveness of AI tools in recognizing phishing attacks. Artificial intelligence technologies have significantly improved the accuracy of phishing detection and supported adaptation to new cyber threats. The analysis shows that the use of AI allows not only to detect attacks in a timely manner, but also to develop prevention strategies. The practical significance of the results lies in the possibility of integrating the developed methods into existing security systems. The work offers a strategic approach combining technological advances and organizational practices to create sustainable information security.

Key words: spam, phishing, information security, machine learning, filtering, methods, protection, challenges.

Авторлар туралы мәліметтер

Аяна Байгабулкызы Какенова* – техника ғылымының магистрі, Л.Н. Гумилев атындағы Еуразия ұлттық университеті, Астана қ.; e-mail: akakenova1001@gmail.com.

Баян Куандыковна Абдураимова – техника ғылымдарының кандидаты, доцент, Л.Н. Гумилев атындағы Еуразия ұлттық университеті, Астана қ.; e-mail: abduraimovabk@mail.ru. ORCID: <https://orcid.org/0000-0003-3913-1895>.

Сая Адильбекқызы Сантеева – Phd, Л.Н. Гумилев атындағы Еуразия ұлттық университеті, Астана қ.; e-mail: saya_santeeva@mail.ru.

Сведения об авторах

Аяна Байгабулкызы Какенова* – магистрант технических наук, Евразийский национальный университет имени Л.Н. Гумилёва, г. Астана; e-mail: akakenova1001@gmail.com.

Баян Куандыковна Абдураимова – кандидат технических наук, доцент, Л.Н. Гумилев атындағы Еуразия ұлттық университеті, г. Астана; e-mail: abduraimovabk@mail.ru. ORCID: <https://orcid.org/0000-0003-3913-1895>.

Сая Адильбекқызы Сантеева – Phd, Евразийский национальный университет имени Л.Н. Гумилёва, г. Астана, e-mail: saya_santeeva@mail.ru.

Information about the authors

Ayana Baygabulkyzy Kakenova* – master of technical sciences, L.N. Gumilyov Eurasian National University, Astana; e-mail: akakenova1001@gmail.com.

Bayan Kuandykovna Abduraimova – candidate of technical sciences, associate professor, L.N. Gumilyov Eurasian National University, Astana; e-mail: abduraimovabk@mail.ru. ORCID: <https://orcid.org/0000-0003-3913-1895>

Saya Adilbekkyzy Santeeva – Phd, L.N. Gumilyov Eurasian National University, Astana; e-mail: saya_santeeva@mail.ru.

Редакцияға енуі 01.11.2024

Өңдеуден кейін түсуі 07.11.2024

Жариялауға қабылданды 08.11.2024