

Мадина Даулетжановна Мусина – аға оқытушы, ақпараттық технологиялар және автоматика кафедрасы М. Дулатов атындағы Қостанай инженерлік-экономикалық университеті, Қазақстан Республикасы; e-mail: madina.madina.musina@mail.ru. ORCID: <https://orcid.org/0009-0004-0610-6938>.

Information about the authors

Dinara Aldasheva* – doctoral student of the department of Software Engineering, A. Baitursynov Kostanay Regional University, Republic of Kazakhstan; e-mail: aldasheva.dinara@mail.ru. ORCID: <https://orcid.org/0009-0000-4990-4308>.

Olga Salykova – candidate of technical sciences, associate professor of the Department of Software Engineering, Kostanay Regional University named after Akhmet Baitursynuly, Republic of Kazakhstan; e-mail: solga0603@mail.ru. ORCID: <https://orcid.org/0000-0002-8681-4552>.

Mikhail Zarubin – candidate of technical sciences, associate professor, department of information technology and automation, Kostanay Engineering and Economics University named after Myrzakyp Dulatov, Republic of Kazakhstan; e-mail: zarubin_mu@mail.ru. ORCID: <https://orcid.org/0000-0002-1415-5244>.

Talgat Zhuaspayev – senior lecturer, department of information technology and automation, Kostanay Engineering and Economics University named after Myrzakyp Dulatov, Republic of Kazakhstan; e-mail: g_talgat_a@mail.ru. ORCID: <https://orcid.org/0009-0000-2240-9729>.

Madina Musina – senior lecturer, department of information technology and automation, Kostanay Engineering and Economics University named after Myrzakyp Dulatov, Republic of Kazakhstan; e-mail: madina.madina.musina@mail.ru. ORCID: <https://orcid.org/0009-0004-0610-6938>.

Поступила в редакцию 05.01.2026

Принята к публикации 16.02.2026

[https://doi.org/10.53360/2788-7995-2026-1\(21\)-7](https://doi.org/10.53360/2788-7995-2026-1(21)-7)

FTAXP: 81.93.29



Б.К. Абдураимова, А.Е. Нұрмұханбетова*

Л.Н. Гумилев атындағы Еуразия ұлттық университеті,
010000 Қазақстан Республикасы, Астана қаласы, Сатпаева көшесі, 2
*e-mail: nuralbina@gmail.com

АҚПАРАТТЫҚ-КОММУНИКАЦИЯЛЫҚ ОРТАДАҒЫ ЗИЯНДЫ БАҒДАРЛАМАЛАРДЫ АНЫҚТАУҒА АРНАЛҒАН ТЕРЕҢ ОҚЫТУҒА НЕГІЗДЕЛГЕН ИНТЕЛЛЕКТУАЛДЫ ГИБРИДТІ МОДЕЛЬДЕР

Аңдатпа: Бұл мақалада ақпараттық-коммуникациялық инфрақұрылымдарда зиянды бағдарламаларды анықтауға арналған терең оқыту әдістеріне негізделген интеллектуалды гибриді модельді зерттеу және әзірлеу ұсынылған. Зерттеудің мақсаты: конволюциялық (CNN), қайталанатын (GRU) және графтық (GNN) нейрондық желілерді онтологиялық ерекшеліктерді қалыпта келтірумен біріктіретін архитектураны құру. Бұл тәсіл зиянды бағдарламалардың статикалық, мінез-құлықтық және құрылымдық сипаттамаларын кешенді талдауға мүмкіндік береді, нәтижесінде күндіз шабуылдар мен полиморфты қауіптерге төзімділікті арттырады. Тәжірибеде тексеру үшін EMBER-2018, Malimg, CICAndMal2017 және CICMalDroid-2020 ашық деректер жиынтығы пайдаланылды. Эксперименттер ұсынылған гибриді архитектураның жіктеу дәлдігін 98.7% және ROC-AUC = 0.99 қамтамасыз ететінін, оқшауланған терең оқыту модельдері мен дәстүрлі машиналық талдау әдістерінен асып түсетінін көрсетті. Онтологиялық білім моделін енгізу жүйенің түсіндірілуін арттырды және жалған оң нәтижелер көрсеткішін 1.7%-ға дейін төмендетті. Әзірленген бағдарламалық жасақтаманың прототипі шешімнің нақты уақыт жүйелерінде (SOC/IDS) қолданылуын көрсетті және практикалық киберқауіпсіздік жүйелерінде одан әрі масштабтау және енгізу әлеуетін көрсетті. Зерттеудің жаңалығы күрделі цифрлық ортада зиянды нысандарды дәлірек және түсінікті түрде анықтауға мүмкіндік беретін мультимодальды талдау мен онтологиялық модельдеуді біріктіруде жатыр.

Түйін сөздер: зиянды бағдарлама; терең оқыту, конволюциялық нейрондық желілер (CNN), қайталанатын нейрондық желілер (GRU), графтық нейрондық желілер (GNN), нәтижесінде күндіз шабуылды анықтау, гибриді архитектура, киберқауіпсіздік.

Кіріспе

Қазіргі таңда зиянды бағдарламалық жасақтаманың жылдам дамуына, әртүрлілігіне және күрделілігіне байланысты зиянды бағдарламаларды анықтау киберқауіпсіздіктегі ең өзекті мәселелердің біріне айналды. Дәстүрлі қолтаңбаға және эвристикалық негізделген әдістер полиморфты, метаморфты және файлсыз шабуылдар сияқты озық қауіптерге қарсы тұру үшін енді жеткіліксіз. Бұлтты есептеулерге, виртуалдандыруға және Заттар интернетіне (IoT) жаһандық ауысу қауіпсіздік ландшафтын одан әрі күрделендірді [18], себебі динамикалық инфрақұрылымыдарда зиянды әрекеттерді дәстүрлі механизмдермен анықтау тиімсіз болып табылады. «Kaspersky Lab» мәліметтері бойынша, 2024 жылы күн сайын орта есеппен 467 000 жаңа зиянды файл анықталған [1]. Бұл 2023 жылмен салыстырғанда 14%-ға өскен [1].

Бұл нәтиже интеллектуалды, автоматтандырылған және масштабталатын анықтау шешімдерінің шұғыл қажеттілігін көрсетеді. Машиналық оқыту (Machine Learning) және терең оқыту (Deep learning) тәсілдері зиянды бағдарламалардың күрделі мінез-құлық және құрылымдық модельдерін зерттеу арқылы айтарлықтай артықшылықтар береді. Бұл гибриді тәсіл арқылы белгісіз немесе нөлдік күн осалдықтарын қауіптерді анықтауға болады. Сонымен қатар, төлем бағдарламаларының, жеткізілім тізбегіне шабуылдардың және мемлекет қаржыландыратын шабуылдардың өсіп келе жатқан экономикалық әсері нақты уақыт режимінде өзгеріп отыратын қауіптерге бейімделе алатын жасанды интеллектке негізделген заманауи киберқауіпсіздік жүйелерін құру қажеттілігін көрсетеді.

Әдебиеттерге шолу. Зиянды бағдарламаларды анықтаудың дәстүрлі әдістері қолтаңба, эвристикалық, статикалық және динамикалық талдауға негізделген. Қолтаңба талдауы алдын ала анықталған үлгілерді (жолдар, хэштер, код фрагменттері) пайдаланады, оларды талданатын файлдың мазмұнымен салыстырады. Бұл тәсіл белгілі қауіптерге қарсы жоғары жылдамдық пен дәлдікпен сипатталады, бірақ жаңа, полиморфты және метаморфты зиянды бағдарламалардың нұсқаларын анықтай алмайды [2]. Кодтағы минималды өзгерістердің өзі, оның ішінде, мысалы, нұсқауларды ауыстыру, бөлім құрылымын өзгерту немесе бумалағыштарды пайдалану қолтаңбаны пайдасыз етеді.

Ал, эвристикалық әдістер күдікті индикаторлар жиынтығын, мысалы, ерекше API шақыруларын, өзін-өзі өзгертетін кодты, жүйелік тізілімдегі өзгерістерді немесе қорғалған каталогтарға кіруді талдайды [3]. Олар қолтаңбаларға қарағанда икемді, бірақ жалған оң нәтижелердің жоғары деңгейде көрсетеді және ережелерді қолмен жаңартуды қажет етеді.

Статикалық талдау қауіпсіз және жылдам, орындалатын файлдарды, олардың метадеректерін, опкодтарын және тәуелділіктерін тексеруге мүмкіндік береді. Дегенмен, ол бағдарламаның нақты әрекетін көрсетпейтіндіктен, шатастыруға және шифрлауға осал [4].

Динамикалық талдау (құм жәшіктері, эмуляторлар, виртуалды машиналар) жүйелік қоңырауларды, желілік белсенділікті және файл операцияларын бақылау үшін бақыланатын ортада нысанды орындайды [5]. Шатастыруға қарсы сенімдірек, бірақ ресурстарды көп қажет етеді және жалтаруға осал, қазіргі заманғы зиянды бағдарлама виртуалдандыруды анықтап, құм жәшігінің іздерін анықтаса, әрекетті өзгерте алады. Беделге негізделген әдістер мен URL сүзгілері белгілі қауіптерге жедел жауап береді, бірақ толығымен жаңартылған дерекқорлар мен бедел серверлеріне тәуелді [6].

Ғылыми мәселе. Қазір зиянды бағдарламаларды анықтау әдістері полиморфты және метаморфты қауіптер санының артуына, шатастыру технологияларын кеңінен қолдануға және шабуылдаушылардың файлсыз шабуылдарға ауысуына байланысты шектеулерге тап болуда. Дәстүрлі қолтаңбаға және эвристикалық тәсілдерге негізделген қажетті бейімделушілік пен жалпылау мүмкіндігін қамтамасыз етпейді, ал кейбір машиналық оқыту әдістері деректердің теңгерімсіздігін, жоғары мінез-құлық өзгергіштігін және репрезентативті мүмкіндіктердің болмауын шеше алмайды. Ғылыми қиындық статикалық, динамикалық және графикалық мінез-құлық ерекшеліктерін біріктіре алатын, гетерогенді инфрақұрылымдарда сенімді және түсіндірілетін зиянды бағдарламаларды анықтауды қамтамасыз ететін гибриді интеллектуалды модельді әзірлеу қажеттілігінде жатыр.

Зерттеу нысаны. Ақпараттық және коммуникациялық ортадағы зиянды бағдарламаларды анықтау және жіктеу процесін жасанды интеллект және терең оқыту әдістерін қолдану.

Зерттеу мақсаты: аралас режимдегі орталарда (Windows, Android, бұлттық және контейнерлік жүйелер) тиімді, сенімді және түсіндірілетін зиянды бағдарламаларды анықтауды қамтамасыз ететін интеллектуалды гибриді терең оқыту моделін әзірлеу.

Зерттеу міндеттері:

- Зиянды бағдарламаларды анықтаудың қолданыстағы әдістеріне аналитикалық шолу жүргізу, олардың күшті және әлсіз жақтарын анықтау және дәстүрлі және бір қабатты терең оқыту модельдерінің шектеулерін анықтау.
- Статикалық, динамикалық және графикалық мүмкіндіктерді бірыңғай нейрондық желіге негізделген архитектурада біріктірудің орындылығын зерттеу.
- Бинарлық кескіндерді талдау үшін CNN, мінез-құлық тізбектері үшін RNN/GRU және графикалық қатынастар үшін GNN біріктіретін гибриді модель архитектурасын әзірлеу.
- Ұсынылған модельдің прототипін енгізу және ашық деректер жиынтықтарында бірқатар эксперименттер жүргізу.
- Ұсынылған модельдің өнімділігін келесі метрикаларды пайдаланып, қолданыстағы шешімдермен салыстыру.
- Модельдің нөлдік күндік шабуылдарға, шатастыруға және файлсыз қауіптерге төзімділігін, сондай-ақ түсіндірілетін жасанды интеллект құралдарын пайдаланып шешімдерді түсіндіру мүмкіндігін бағалау.
- Ұсынылған тәсілді ақпараттық-коммуникациялық инфрақұрылымға арналған практикалық киберқауіпсіздік жүйелерінде одан әрі дамыту және енгізу салаларын анықтау.

Әдіснама

Бұл бөлімде зерттеу әдіснамасы, соның ішінде пайдаланылған материалдардың сипаттамалары, зерттеу сұрақтары мен гипотезалардың тұжырымдамасы, жұмыс кезеңдері және қойылған мақсатқа жету үшін қолданылатын әдістер сипатталған.

Зерттеу материалының сипаттамалары

Зерттеу материалы ретінде Windows және Android платформаларын қамтитын ашық деректер жиынтығын пайдаланылады, себебі бұл гибриді және кросс-платформалы зиянды бағдарламаларды анықтау модельдерін әзірлеуге мүмкіндік береді.

Windows орындалатын файлдарын (Portable Executable, PE форматы) талдау үшін деректер жиынтығында шамамен 1,1 миллион үлгі (900 000 оқыту және 200 000 тест) қамтитын және мүмкіндіктер мен сынып белгілері көрсетілген EMBER-2018 деректер жиынтығы қарастырылады [6]. Бинарлы деректерді визуализациялау үшін 25 отбасынан 9339 зиянды бағдарлама даналарын қамтитын, суреттер түрінде ұсынылған Malimg деректер жиынтығы пайдаланылды және бұл деректерді пайдалану жіктеу тапсырмасына компьютерлік көру (CNN) әдістерін қолдануға мүмкіндік береді [7].

Android мобильді ортасы үшін зерттеу Канада киберқауіпсіздік институты ұсынған CICAndMal2017 және CICMalDroid-2020 деректер жиынтығына негізделген. Біріншісі 10 854 қосымшаны (оның ішінде 4 354 зиянды бағдарлама), екіншісінде статикалық және динамикалық ерекшеліктері бар (API қоңыраулары, желілік ағындар, жүйелік қоңыраулар) 17 000-нан астам дананы қамтиды [8].

Сонымен қатар, зерттеуде API тәуелділіктерін график түрінде көрсету әдістері (басқару ағыны және шақыру графигі) қолданылады, бұл мінез-құлық үлгілерін талдау үшін графикалық нейрондық желілерді (GNN) құруға мүмкіндік береді [9]. Деректердің семантикалық үйлесімділігін жақсарту үшін шабуыл түрлері, мүмкіндіктері және зиянды бағдарламалардың әрекеті арасындағы қатынастарды сипаттайтын OWL форматындағы онтологиялық білім моделі әзірленді [10].

Осылайша, материалдар жиынтығы талдаудың үш басым әдісін – статикалық, динамикалық және құрылымдық-графикалық тәсілдерін қамтиды, бұл нәтижелердің кешенділігі мен қайталануын қамтамасыз етеді.

Зерттеу сұрақтары және гипотеза

Зерттеу сұрақтары:

- Статикалық, динамикалық және графикалық мүмкіндіктерді біріктіретін мультимодальды гибриді архитектура жеке бір қабатты модельдермен (тек CNN немесе тек RNN) салыстырғанда нөлдік күндік шабуылдарға дәлдікті және төзімділікті жақсартта ма?
- Онтологиялық ерекшеліктерді қалыпқа келтіруді пайдалану әртүрлі деректер жиынтығы мен платформалар арасында берілген кезде модельдердің жалпылау мүмкіндігін жақсартта ма?
- Түсіндірілетін жасанды интеллект әдістерін (Grad-CAM, GNNExplainer) біріктіру зиянды үлгілерді анықтаудың жоғары деңгейін сақтай отырып, шешімдердің түсіндірілуін жақсартта ала ма?

Ұсынылған гипотеза (тезис): Зерттеудің гипотезасы бойынша, онтологиялық тұрғыдан нормаланған мүмкіндіктерді мультимодальды терең оқыту архитектурасымен (CNN + RNN/GRU + GNN) біріктіру оқшауланған модельдермен салыстырғанда жіктеу сапасының статистикалық тұрғыдан айтарлықтай жақсаруын ($F1\text{-score} > 0,95$ және $ROC\text{-AUC} > 0,98$ метрикаларына сәйкес) және обфускация мен нөлдік күндік шабуылдарға төзімділікті қамтамасыз етеді [11, 12].

Бірінші кезең зиянды бағдарламаларды анықтауға арналған ғылыми басылымдар мен практикалық шешімдерді аналитикалық шолуды қамтиды, себебі бұл қолтаңбаға негізделген әдістерден бастап терең нейрондық желілерге негізделген гибриді модельдерге дейін негізгі даму бағыттарын анықтауға мүмкіндік берді [2]. Өртүрлі нейрондық желі архитектураларының тиімділігіне, мүмкіндіктерді көрсету әдістеріне және модельдердің полиморфты және файлсыз шабуылдарға төзімділігіне қатысты зерттеу сұрақтары тұжырымдалды.

Екінші кезең деректерді жүйелеуді және тазартуды, соның ішінде сынып теңгерімсіздігін жою үшін қайталануды алып тастауды, қалыпқа келтіруді және кеңейтуді қамтиды. Функция кеңістігін құру үшін статикалық, динамикалық және графиктік сипаттамалар пайдаланылды. Статикалық мүмкіндіктер PE/DEX метадеректерін, opcode тізбектерін және кітапхана импорттарын қамтиды. Ал, динамикалық мүмкіндіктер жүйелік шақырулар мен желілік ағындарды, ал графиктік мүмкіндіктер API шақыру графиктері арқылы көрсетілген функциялар мен жадтағы нысандар арасындағы қатынастарды қамтиды [4].

Үшінші кезең зиянды бағдарламалар отбасылары, мүмкіндіктер, шабуыл әдістері және инфрақұрылым контексті арасындағы қатынастарды сипаттайтын OWL онтологиясын жасауды қамтиды. Онтология деректер семантикасын біріктіру және машиналық оқыту нәтижелерін түсіндіру үшін қолданылады [13].

Төртінші кезең көрнекі екілік кескіндерді талдау үшін конволюциялық нейрондық желілерді (CNN), жүйелік шақыру тізбектерін талдау үшін рекуррентті нейрондық желілерді (RNN, GRU) және топологиялық қатынастарды көрсету үшін графтық нейрондық желілерді (GNN) біріктіретін гибриді архитектураны құруға және оқытуға арналды. Модальділіктер толық байланысты қабатпен кеш біріктіруді қолдану арқылы біріктірілді. Оқыту үшін Adam және RMSProp оңтайластырғыштары, сондай-ақ кросс-энтропия және фокальды жоғалту функциялары пайдаланылды.

Бесінші кезеңде әзірленген модельдер келесі көрсеткіштерді қолдана отырып, сынақ жиынтықтарында сынақтан өткізіліп, оңтайландырылды: дәлдік (Accuracy, Precision), еске түсіру (Recall), $F1\text{-score}$, $ROC\text{-AUC}$ және $FPR@TPR=95\%$. Модельдердің нөлдік күндік шабуылдарға және жасырын зиянды әрекеттерге төзімділігіне ерекше назар аударылды. Осы мақсатта оқыту жиынынан тұтас отбасыларды алып тастап (отбасы аралық бөлу) эксперименттер жүргізілді, бұл бізге модельдің білімді жалпылау қабілетін бағалауға мүмкіндік берді [8].

Соңғы кезеңде деректерді жинау модулін, машиналық оқыту модулін және нәтижелерді визуализациялау интерфейсін қамтитын бағдарламалық платформа енгізіледі. Бағдарламалық жасақтаманы енгізу TensorFlow, PyTorch, Scikit-learn және Keras кітапханаларын пайдаланып Python тілінде жазылған. Нәтижелер гибриді архитектураның дәлдігі жоғары екенін (классикалық CNN немесе RNN модельдеріне қарағанда 6-12% жоғары) және ұқсас шешімдермен салыстырғанда жалған оң нәтижелердің төмен деңгейін көрсететінін көрсетті [14].

Осылайша, ұсынылған әдіснама аналитикалық, эксперименттік және инженерлік тәсілдерді біріктіреді және іс жүзінде маңызды нәтижеге қол жеткізуге бағытталған: бұлтқа негізделген инфрақұрылымды қорғауға және ақпараттық жүйелердің кибертөзімділігін қамтамасыз етуге қолданылатын зиянды бағдарламаларды интеллектуалды талдауға арналған бейімделгіш жүйені құру.

Зерттеу әдістері

Бұл зерттеудің әдіснамасы ақпараттық-коммуникациялық инфрақұрылымдардағы интеллектуалды гибриді зиянды бағдарламаларды анықтау модельдерінің тиімділігін теориялық талдау, инженерлік модельдеу және эксперименттік тестілеудің үйлесіміне негізделген. Бұл тәсіл бұрын белгісіз және полиморфты қауіптерді анықтай алатын, нөлдік күндік шабуылдарға төзімді және нақты әлемдегі корпоративтік және бұлттық орталарға қолданылатын жүйені құруға бағытталған.

Зерттеуде әртүрлі зиянды бағдарламалар түрлері мен платформаларын қамтитын ашық деректер жиынтығы пайдаланылды. Windows орындалатын файлдарын талдау үшін аннотацияланған мүмкіндіктері бар миллионнан астам үлгілерді қамтитын EMBER-2018 деректер жиынтығы пайдаланылды. Ал, файлдарды суреттер ретінде жобалау үшін екілік құрылымдарды визуализациялау қолданылды. Android мобильді платформасы үшін статикалық және динамикалық мүмкіндіктерді (API шақырулары, желілік ағындар, жүйелік шақырулар) қамтитын CICAndMal2017 және CICMalDroid-2020 пайдаланылды. Бұл деректер көздері тәжірбиелердің әртүрлілігі мен кешенділігін қамтамасыз етті, бұл бізге әртүрлі домендер арасында ауысқан кезде модельдердің сенімділігін бағалауға мүмкіндік берді.

Зерттеу нәтижелері

Зерттеу нәтижелері ақпараттық-коммуникациялық инфрақұрылымдарда зиянды бағдарламаларды анықтау үшін конволюциялық, қайталанатын және графикалық нейрондық желілерді біріктіруді ұсынатын интеллектуалды гибриді архитектураның тиімділігін көрсетеді. Тәжірбиелер әртүрлі деректер түрлерін (статикалық, динамикалық және құрылымдық) талдауды, қолданыстағы әдістермен салыстырмалы бағалауды және нақты ақпараттық қауіпсіздік жүйелеріне ұқсас жағдайларда модельдерді тестілеуді қамтыды.

Бірінші кезеңде дәстүрлі машиналық оқыту алгоритмдерін біріктіретін базалық модельдер SVM, Random Forest және LightGBM оқытылып, сынақтан өткізілді. Олардың нәтижелері классикалық тәсілдердің белгілі шектеулерін растады, яғни белгілі үлгілер үшін жоғары дәлдікті көрсетсе, жаңа және бұрын көрінбеген мысалдарға жалпылау мүмкіндігі төмен болып табылады. Бұл модельдер үшін орташа F1 – ұпайлары (F1 – score) 0.91-ден аспады, ал ROC-AUC мәндері 0.89-0.93 деңгейінде қалды, бұл полиморфты және файлсыз шабуылдарға шектеулі бейімделуді көрсетеді [6, 8].

Екіншіден, тәжірбиелер жиынтығында терең оқыту модельдерінің CNN, GRU және GNN өнімділігі бағаланды, әрқайсысы өз модальдылығының деректерін талдады. Malimg деректер жиынтығынан визуализацияланған екілік файлдарда оқытылған CNN моделі 95,1% дәлдік пен F1-ұпайы (F1-score) 0,94 көрсетті, бұл компьютерлік көру әдістерінің екілік құрылымдарды талдау міндетіне қолданылуын растады. CICMalDroid-2020 жүйелік шақыру тізбектерімен жұмыс істейтін GRU моделі F1 ұпайы 0,93 және бағдарламаның мінез-құлқындағы жасырын уақытша үлгілерді анықтау мүмкіндігін көрсетті. API тәуелділік графиктерін талдайтын GNN 96,2% дәлдікке қол жеткізді және контекстік және топологиялық қатынастарды түсіру арқылы жаңа зиянды бағдарламалар нұсқаларына төзімділігін растады [7, 9].

Ең жоғары нәтижелерге үш тармақтың барлығын біріктіру механизмі бар бірыңғай гибриді CNN+GRU+GNN архитектурасы арқылы қол жеткізіледі. Соңғы дәлдік 98.7%, F1-ұпайы (F1-score) 0.98, ал ROC-AUC 0.99 болды. API шақыруларын MITRE ATT&CK тактикасы мен әдістерімен байланыстыратын онтологиялық ерекшеліктерді қалыпқа келтіруді қосу арқылы жалған оң нәтиже көрсеткіші (FPR) 95% , ал TPR кезінде 1.7%-ға дейін төмендеді [15]. Бұл онтологиялық білім моделін пайдалану жіктеуіштің түсіндірілуін және жалпылау қабілетін жақсартатынын көрсетеді.

Модельдердің салыстырмалы нәтижелері 1-кестеде келтірілген. Онтологияны қолдайтын гибриді модель тек классикалық алгоритмдерден ғана емес, сонымен қатар MAD-NET моделін [8] қоса алғанда, ең жақсы гибриді шешімдерден де асып түсетінін, ұсынылған шешімдердің жақсы түсіндірілуін және дәлдіктің ұқсас деңгейінде есептеу күрделілігінің төмен екенін көрсетеді.

Қолданылған тәжірбиелер бізге әрбір әдістің үлесін сандық бағалауға мүмкіндік берді. CNN-ді алып тастау дәлдікті 5.4%-ға төмендетті, бұл екілік файлдардың құрылымдық үлгілерін талдаудың маңыздылығын көрсетті. GRU-ларды алып тастау F1 – ұпайын (F1 – score) шамамен 6%-ға төмендетті, бұл уақытша тәуелділіктердің рөлін айқындайды. GNN-дерді алып тастау ROC-AUC-ті 3-4%-ға төмендетті, бұл кодты өзгертуге беріктік үшін графикалық байланыс талдауының маңыздылығын растайды [11]. Осылайша, гибриді модель синергетикалық әсер көрсетті, үш тәсілді біріктіріп қолдану олардың жеке әсерлерінің қосындысынан да жоғары нәтиже берді.

Кесте 1 – Зиянды бағдарламаларды анықтау модельдерінің салыстырмалы талдауы

№	Модельдің түрі	Пайдаланылған деректер	Дәлдік (%)	F1-ұпайы	ROC-AUC	FP-rate (%)
1	SVM	EMBER-2018	89.7	0.88	0.91	7.4
2	Random Forest	EMBER-2018, CICAndMal2017	91.2	0.90	0.93	6.1
3	LightGBM	EMBER-2018	93.5	0.92	0.95	5.2
4	CNN	Malimg	95.1	0.94	0.97	4.3
5	GRU	CICMalDroid-2020	94.6	0.93	0.96	4.6
6	GNN	CFG/API-графтары	96.2	0.95	0.97	3.9
7	CNN+GRU	Визуализация + әрекеттері	97.4	0.96	0.98	2.8
8	CNN+GRU+GNN	Барлық әдістер, кеш біріктіру	98.7	0.98	0.99	2.1

Шешімдердің түсіндірілуін талдау үшін Grad-CAM [16] және GNNExplainer [17] қолданылды. Grad-CAM бізге жіктеуге ең көп әсер ететін екілік кескіндердің аймақтарын визуализациялауға мүмкіндік берді, ал GNNExplainer API граф түйіндері арасындағы негізгі байланыстарды анықтады. Бұл тәсілдер жүйенің жұмысының ашықтығын және ақпараттық қауіпсіздік мамандарының сенімін арттырады.

2-кестеден көрініп тұрғандай, LightGBM негізіндегі модельдер [19] F1-score = 0.9547 деңгейін көрсетсе, Random Forest негізіндегі механизмдер [21] 98.13% дәлдікке жеткен. NLP және мінез-құлықтық белгілерге негізделген модельдер [22] салыстырмалы түрде төмен нәтиже ($F1 \approx 0.858$) көрсетті.

Кесте 2 – 2023–2025 жылдардағы зерттеулермен салыстыру

№	Автор (жыл)	Деректер жиыны	Accuracy (%)	Recall	F1-score
1	AlOmari et al. [19], 2023	CICMalDroid2020	95.49	0.947	0.9547
2	Azeem et al. [20], 2024	UNSW-NB15	98.13	0.9998	0.9998
3	Bakshi et al. [21], 2025	Figshare datasets	89.30	0.9821	0.9821
4	Gupta et al. [22], 2025	Behavioral sandbox logs	99.98	0.871	0.858
5	Ұсынылған модель, 2025	CNN+GRU+GNN	98.7	0.98	0.98

Ұсынылған CNN+GRU+GNN гибриді архитектурасы 98.7% дәлдік, $F1=0.98$ және $ROC-AUC=0.99$ нәтижелерін көрсетті, бұл бірмодальді немесе тек ансамбльдік әдістермен салыстырғанда жоғары немесе бәсекеге қабілетті нәтиже екенін дәлелдейді.

Айта кету керек, әртүрлі зерттеулерде қолданылған деректер жиындары мен бағалау протоколдары әртүрлі болғандықтан, нәтижелер тікелей салыстыру үшін толық эквивалентті емес. Дегенмен, алынған көрсеткіштер ұсынылған модельдің қазіргі заманғы әдістерімен тең немесе жоғары деңгейде екенін көрсетеді.

Зерттеу нәтижелері

Эксперименттік зерттеулер зиянды бағдарламаларды анықтау үшін конволюциялық (CNN), рекуррентті (GRU) және графтық (GNN) нейрондық желілерді біріктіретін ұсынылған гибриді архитектураның жоғары тиімділігін растады.

Негізгі машиналық оқыту әдістері (SVM, Random Forest, LightGBM) алдыңғы зерттеулерге сәйкес 89-93% дәлдік және ROC-AUC 0.95 дейін көрсетті. Терең оқыту модельдері жоғары өнімділікке қол жеткізді: екілік кескіндердегі CNN (Malimg) 95.1% дәлдікке, мінез-құлық деректеріндегі GRU (CICMalDroid-2020) 94.6% және API шақыру графиктеріндегі GNN 96.2% дәлдікке қол жеткізді.

Ұсынылған CNN+GRU+GNN гибриді моделі оқшауланған модельдерден 6-12% артық нәтижесімен 98.7%, F1-ұпай ($F1-score$) 0.98 және $ROC-AUC=0.99$ дәлдігін көрсетті. Ерекшеліктерді онтологиялық қалыпқа келтіруді енгізу жалған оң нәтижелер санын 1.7%-ға дейін, ал $TPR=95\%$ төмендетуге мүмкіндік берді.

Ұсынылған модель жалған оң позитивтердің төмен деңгейін көрсетсе де ($FP-rate = 1.7\%$), SOC/IDS ортасында практикалық қолдану үшін тіпті мұндай көрсеткіш

қауіпсіздік талдаушыларына айтарлықтай жүктеме әкелуі мүмкін. Сондықтан қате жіктеу жағдайларына қосымша талдау жасалды.

FP болжамдарын талдау модельдің келесі жағдайларда жиі қателесетінін көрсетті:

1. Қапталған зиянсыз орындалатын файлдар. Қаптағыштарды пайдаланатын кейбір заңды бағдарламалар (мысалы, UPX) зиянды үлгілерге ұқсас энтропия сипаттамаларына ие. Көрнекі CNN модулі мұндай құрылымдарды күдікті деп түсіндіреді.

2. Қарқынды API шақырулары бар қолданбалар. Заңды әкімшілік құралдары мен жүйелік утилиталары трояндарға ұқсас мінез-құлық үлгілерін көрсетеді (жиі желілік қосылымдар, файлдық жүйеге кіру), бұл зиянды бағдарлама класының ықтималдығының жоғарылауына әкеледі.

3. CFG графиктерінің жоғары құрылымдық байланысы. GNN модулі кейде күрделі бағдарламалық жасақтама архитектураларын графикалық қосылымдардың жоғары тығыздығына байланысты зиянды деп жіктейді, бұл түсініксіз, бірақ заңды бағдарламаларға да тән.

4. Нақты әлемдегі ортадағы кластардың теңгерімсіздігі. SOC ортасында зиянсыз трафиктің үлесі зиянды бағдарламадан айтарлықтай асып түседі. Тіпті жоғары ерекшелікпен де, ықтималдық шегіндегі шағын қателік FP абсолютті санының артуына әкелуі мүмкін.

Жалған оң талдау

Жүйе күніне 1 000 000 оқиғаны талдайды деп есептейік. FP көрсеткіші 1,7% болғанда, бұл айтарлықтай операциялық жүктемені тудырады.

$$FP = 0.017 \times 1\,000\,000 = 17\,000$$

Дегенмен, жіктеу шегін оңтайландыру (шекті реттеу) және ROC қисығын пайдалану арқылы TPR > 93% сақтай отырып, FP-ны 0,8-1,0%-ға дейін төмендетуге болады.

Осылайша, жалпы дәлдіктің жоғары деңгейіне (98,7%) қарамастан, қателіктерді талдау жалған оң нәтижелердің негізінен құрылымдық тұрғыдан күрделі немесе пакеттелген заңды бағдарламалар жағдайында болатынын көрсетеді. Бұл модельді нақты әлемдегі SOC инфрақұрылымдарында енгізу кезінде бейімделгіш шекті баптау және контекстік сүзгілеумен гибриді тәсілдің қажеттілігін растайды.

Ғылыми нәтижелерді талқылау

Зерттеу нәтижелері әртүрлі деректер модальділіктерін статикалық, динамикалық және графикалық біріктіру жіктеу сапасын және нөлдік күндік шабуылдарға төзімділікті айтарлықтай жақсартатынын растайды. Бұл тұжырымдар Shaukat, Zhu және тағы басқа ғалымдардың зерттеулерінің [15, 9] нәтижелерімен сәйкес келді, олар мультимодальды және онтологияға негізделген тәсілдердің тиімділігін атап өтті.

Гибриді модель синергетикалық әсерді көрсетеді: әрбір тармақ (CNN, GRU, GNN) жалпы дәлдікке ерекше үлес қосады. Түсіндірілетін жасанды интеллект әдістерін (Grad-CAM, GNNExplainer) қолдану шешімдердің түсіндірілуін және жүйені автоматтандырылған киберқауіпсіздік жүйелерінде практикалық енгізудің орындылығын растады.

Осылайша, әзірленген архитектура жоғары дәлдік, жылдамдық және шешім қабылдау ашықтығы арасындағы тепе-теңдікті қамтамасыз етеді, бұл инфрақұрылымдардағы зиянды шабуылдарды бақылау және алдын алу үшін перспективті құралға айналдырады.

Қорытынды

Бұл зерттеу ақпараттық және коммуникациялық инфрақұрылымдардағы зиянды бағдарламаларды анықтауға арналған терең оқыту әдістеріне негізделген интеллектуалды гибриді архитектураны әзірледі, енгізді және эксперименттік түрде растады. Ұсынылған модель үш қосымша аналитикалық тәсілді біріктіреді: статикалық екілік мүмкіндіктерді талдауға арналған конволюциялық нейрондық желілер (CNN), мінез-құлық тізбектерін өңдеуге арналған қайталанатын бірлік желілері (GRU) және топологиялық тәуелділіктерді талдауға арналған графтық нейрондық желілер (GNN). Бұл комбинация жіктеу дәлдігін, нөлдік күндік шабуылдарға төзімділікті және шешімдердің түсіндірілуін жақсартты.

EMBER-2018, Malimg, CICAndMal2017 және CICMalDroid-2020 ашық деректер жиынтықтарында жүргізілген эксперименттер гибриді модельдің дәстүрлі және бір деңгейлі архитектуралардан асып түсетінін көрсетті. Атап айтқанда, жіктеу дәлдігі 98.7%-ға жетті, ал ROC-AUC 0.99 болды, бұл оқшауланған модельдерге қарағанда 6-12%-ға жоғары [2], [3], [5]. Онтологиялық ерекшеліктерді қалыпқа келтіруді қосу модельдің жаңа және полиморфты үлгілерге беріктігін арттырды, жалған оң нәтиже көрсеткішін 95% TPR кезінде 1.7%-ға дейін төмендетті [6].

Бұл зерттеудің ғылыми жаңалығы онтологиялық тұрғыдан байытылған ерекшеліктер мен мультимодальды талдауды біріктіруде жатыр, бұл зиянды кодтың құрылымдық, мінез-құлықтық және графикалық аспектілері арасындағы семантикалық үйлесімділікті қамтамасыз етеді. Қолданыстағы шешімдерден айырмашылығы, ұсынылған модель бір ғана деректер санатымен шектелмейді, керісінше оларды бағдарлама әрекеттері, функциялары және желілік оқиғалар арасындағы байланыстарды тәуелсіз анықтауға қабілетті ортақ көрініс жүйесіне біріктіреді.

Бұл жұмыстың практикалық маңыздылығы әзірленген архитектураны нақты уақыт режиміндегі киберқауіптерді бақылау және талдау жүйелерінің прототиптерін құру үшін пайдалануға болатындығында. Орташа қорытынды уақыты 35-40 мс болатын бұл модельді өнімділіктің төмендеуінсіз SOC және IDS инфрақұрылымдарында пайдалануға мүмкіндік береді деп болжануда. Explainable AI (Grad-CAM, GNNExplainer) пайдалану шешім қабылдаудың ашықтығын қамтамасыз етеді, бұл кейінгі оқиғаларды талдау және қауіпсіздік жүйелерінің аудиті үшін маңызды.

Бұл нәтижелер терең оқыту мен онтологиялық модельдеудің үйлесімі автоматтандырылған зиянды бағдарламаларды анықтауда жаңа бағыт қалыптастырады деген болжамды растайды. Әрі қарай зерттеудің әрі қарай дамуына деректер жиынтығын кеңейту, архитектураны басқа қауіп түрлеріне (файлсыз шабуылдар мен IoT зиянды бағдарламаларын қоса алғанда) бейімдеу және ұсынылған жүйені бұлттық платформалар мен болжамды жауап беру жүйелеріне біріктіру кіреді.

Осылайша, ұсынылған әдіснама және алынған нәтижелер киберқауіпсіздікке ғылыми және қолданбалы тәсілдерді дамытуға айтарлықтай үлес қосады, масштабталатын ақпараттық және коммуникациялық орталарда интеллектуалды, бейімделгіш зиянды бағдарламаларды талдау жүйелерін құру мүмкіндіктерін ашады.

Әдебиеттер тізімі

1. CHislo – goda resheniya laboratorii Laboratoriya Kasperskogo ezhednevno obnaruzhivayut 467 tysyach novyh vredonosnyh fajlov // Press-reliz. – Moskovskaya obl., 4 dek. 2024. <https://www.kaspersky.ru/about/press-releases/chislo-goda-resheniya-laboratorii-kasperskogo-ezhednevno-obnaruzhivayut-467-tysyach-novyh-vredonosnyh-fajlov>.
2. Malware visualization for deep learning detection / Q. Cui et al // IEEE Access. – 2022. – T. 10. – P. 24435-24447.
3. Saxe J. Deep Neural Network Based Malware Detection Using Two-Dimensional Binary Program Features / J. Saxe, K. Berlin // Proc. 10th Int. Conf. on Malicious and Unwanted Software (MALWARE). – 2015. – P. 11-20. <https://doi.org/10.1109/MALWARE.2015.7413680>.
4. Hybrid static and dynamic analysis for malware detection / A. Marastoni et al // Computers & Security. – 2024. – Vol. 135. – Article 109934.
5. AV-Comparatives. Real-World Protection Test 2024 – Summary Report // AV-Comparatives.org. – 2024.
6. Anderson H.S. EMBER: An open dataset for training static PE malware machine learning models / H.S. Anderson, P. Roth // arXiv. – 2018. – arXiv:1804.04637.
7. Malware Images: Visualization and Automatic Classification / L. Nataraj et al // Proc. International Symposium on Visualization for Cyber Security (VizSec). – Jul 2011. – P. 1-7. <https://doi.org/10.1145/2016904.2016908>.
8. Poornima P. Automated malware detection using machine learning and deep learning approaches for android applications / P. Poornima, G. Mahalakshmi // Expert Systems with Applications. – 2024. – T. 237. – Article 121456.
9. Graph Neural Networks for Malware Detection: A Survey / Y. Zhu et al // IEEE Access. – 2023. – T. 11. – P. 85123-85147.
10. Ding Y. Ontology-based knowledge representation for malware individuals and families / Y. Ding, R. Wu, X. Xiao // Computers & Security. – 2019. – Vol. 87. – Art. 101574. <https://doi.org/10.1016/j.cose.2019.101574>.
11. Deldar F. Deep Learning for Zero-day Malware Detection and Classification: A Survey / F. Deldar, M. Abadi // ACM Computing Surveys. – 2023. – Vol. 56, № 2. – P. 1-37. <https://doi.org/10.1145/3605775>.

12. ScaleMalNet: A Scalable Hybrid Deep Learning Model for Malware Classification / R. Vinayakumar et al // Future Generation Computer Systems. – 2020. – Т. 115. – P. 280-292.
13. Alshoulie M. Deep Learning Approaches for Malware Detection: A Comprehensive Review of Techniques, Challenges, and Future Directions / M. Alshoulie, A. Mehmood // IEEE Access. – 2025. – Т. 13. – P. 118652-118677. <https://doi.org/10.1109/ACCESS.2025.3582875>.
14. Machine learning based fileless malware traffic classification using image visualization / F.A. Demmese et al // Cybersecurity. – 2023. – Т. 6, № 1. – Art. 32. <https://doi.org/10.1186/s42400-023-00170-z>.
15. A framework for detecting zero-day exploits in network flows / A. Touré et al // Computer Networks. – 2024. – Vol. 224. – Art. 110476. <https://doi.org/10.1016/j.comnet.2024.110476>.
16. Grad-CAM: Visual Explanations from Deep Networks via Gradient-Based Localization / R.R. Selvaraju et al // Proceedings of the IEEE International Conference on Computer Vision (ICCV). – 2017. – P. 618-626. <https://doi.org/10.1109/ICCV.2017.74>.
17. GNNExplainer: Generating Explanations for Graph Neural Networks / R. Ying et al // Advances in Neural Information Processing Systems (NeurIPS). – 2019. – Vol. 32. – P. 9240-9251. <https://doi.org/10.48550/arXiv.1903.03894>.
18. Abduraimova B. Comparative study of machine learning applications in malware forensics / B. Abduraimova, S. Gnatyuk, A. Nurmukhanbetova // Proceedings of the Workshop on Cybersecurity Providing in Information and Telecommunication Systems II (CPITS-II 2024). – Kyiv, Ukraine, October 26, 2024. – CEUR Workshop Proceedings, Vol. 3826. – P. 139-152. <https://ceur-ws.org/Vol-3826/paper13.pdf> (date of request: 03.11.2025).
19. AlOmari H. A comparative analysis of machine learning algorithms for Android malware detection / H. AlOmari, Q.M. Yaseen, M.A. Al-Betar // Procedia Computer Science. – 2023. – Vol. 220. – P. 763-768. <https://doi.org/10.1016/j.procs.2023.03.101>.
20. Analyzing and comparing the effectiveness of malware detection: A study of machine learning approaches / M. Azeem et al // Heliyon Computer Science. – 2024. – Vol. 10, № 1. – Art. e24058. <https://www.sciencedirect.com/science/article/pii/S2405844023107821>.
21. A robust machine learning-based mechanism for malware attack detection and analysis / R. Bakshi et al // Procedia Computer Science. – 2025. – Vol. 242. – P. 876-885. <https://www.sciencedirect.com/science/article/pii/S1877050925010646>.
22. Efficient malware detection using NLP and deep learning model / U. Gupta et al // Alexandria Engineering Journal. – 2025. – Vol. 124. – P. 550-564. <https://www.sciencedirect.com/science/article/pii/S1110016825004260>.

Б.К. Абдураимова, А.Е. Нурмуханбетова*

Евразийский национальный университет им. Л.Н. Гумилева,
010000 Республика Казахстан, г. Астана, ул. Сатпаева, 2
*e-mail: nuralbina@gmail.com

ИНТЕЛЛЕКТУАЛЬНЫЕ ГИБРИДНЫЕ МОДЕЛИ НА ОСНОВЕ ГЛУБОКОГО ОБУЧЕНИЯ ДЛЯ ОБНАРУЖЕНИЯ ВРЕДНОСНЫХ ПРОГРАММ В ИНФОРМАЦИОННО-КОММУНИКАЦИОННОЙ СРЕДЕ

В данной работе представлено исследование и разработка интеллектуальной гибридной модели на основе методов глубокого обучения для обнаружения вредоносного программного обеспечения (ВПО) в информационно-коммуникационных инфраструктурах. Цель исследования заключается в создании архитектуры, объединяющей сверточные (CNN), рекуррентные (GRU) и графовые (GNN) нейронные сети с онтологической нормализацией признаков. Такой подход обеспечивает комплексный анализ статических, поведенческих и структурных характеристик вредоносных программ, повышая устойчивость к атакам нулевого дня и полиморфным угрозам. Для экспериментальной проверки использовались открытые наборы данных EMBER-2018, Malimg, CICAndMal2017 и CICMalDroid-2020. Проведённые эксперименты показали, что предложенная гибридная архитектура обеспечивает точность классификации 98.7% и ROC-AUC 0.99, превосходя изолированные модели глубокого обучения и традиционные методы машинного анализа. Введение онтологической модели знаний позволило повысить интерпретируемость системы и снизить уровень ложноположительных срабатываний до 1.7 %. Разработанный программный прототип продемонстрировал применимость решения в системах реального времени (SOC/IDS) и показал потенциал для дальнейшего масштабирования и внедрения в практические системы

кибербезопасности. Научная новизна исследования заключается в интеграции мультимодального анализа и онтологического моделирования, что обеспечивает более точное и объяснимое выявление вредоносных объектов в сложных цифровых средах.

Ключевые слова: вредоносное программное обеспечение, глубокое обучение, сверточные нейронные сети (CNN), рекуррентные нейронные сети (GRU), графовые нейронные сети (GNN), обнаружение атак нулевого дня, гибридная архитектура, кибербезопасность.

B.K. Abduraimova, A.E. Nurmukhanbetova*

L.N. Gumilyov Eurasian National University,
010000, Republic of Kazakhstan, Astana, Satpayev Street
*e-mail: nuralbina@gmail.com

INTELLIGENT DEEP LEARNING-BASED HYBRID MODELS FOR MALWARE DETECTION IN THE ICT ENVIRONMENT

This paper presents the research and development of an intelligent hybrid model based on deep learning methods for malware detection in information and communication infrastructures. The goal of the study is to create an architecture combining convolutional (CNN), recurrent (GRU), and graph (GNN) neural networks with ontology-based feature normalization. This approach provides a comprehensive analysis of the static, behavioral, and structural characteristics of malware, increasing resilience to zero-day attacks and polymorphic threats. The EMBER-2018, Malimg, CICAndMal2017, and CICMalDroid-2020 open datasets were used for experimental validation. The experiments showed that the proposed hybrid architecture provides a classification accuracy of 98.7% and a ROC-AUC of 0.99, outperforming isolated deep learning models and traditional machine analysis methods. The introduction of an ontological knowledge model increased the system's interpretability and reduced the false-positive rate to 1.7%. The developed software prototype demonstrated the solution's applicability in real-time systems (SOC/IDS) and demonstrated potential for further scaling and implementation in practical cybersecurity systems. The scientific novelty of the study lies in the integration of multimodal analysis and ontological modeling, which enables more accurate and explainable detection of malicious objects in complex digital environments.

Key words: malware, deep learning, convolutional neural networks (CNN), recurrent neural networks (GRU), graph neural networks (GNN), zero-day attack detection, hybrid architecture, cybersecurity.

Авторлар туралы мәліметтер

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>.

Альбина Ерланқызы Нұрмұханбетова* – докторант кафедры информационной безопасности факультета информационных технологий ЕНУ им. Л.Н. Гумилева, Астана, Казахстан; e-mail: nuralbina9898@gmail.com.

Сведения об авторах

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>.

Альбина Ерланқызы Нұрмұханбетова* – Л.Н. Гумилева ат. ЕҰУ Ақпараттық технологиялар факультетінің Ақпараттық қауіпсіздік кафедрасының докторанты, Астана, Қазақстан; e-mail: nuralbina9898@gmail.com.

Information about the authors

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>.

Albina Erlankyzy Nurmukhanbetova* – doctoral candidate in the Department of Information Security, Faculty of Information Technologies, L.N. Gumilyov Eurasian National University, Astana, Kazakhstan; e-mail: nuralbina9898@gmail.com.

Редакцияға енуі 05.01.2026

Өңдеуден кейін түсуі 16.02.2026

Жариялауға қабылданды 17.02.2026