

КОГНІТИВНИЙ РІВЕНЬ БЕЗПЕКИ ЯК НАДБУДОВА ПРИКЛАДНОГО РІВНЯ OSI: АНАЛІТИЧНА МОДЕЛЬ, АРХІТЕКТУРА ТА ПЕРСПЕКТИВИ РОЗВИТКУ

Д.І. Прокопович-Ткаченко^{1,2,3}, О.В. Черкаський³, Д.О. Черкаський⁴,
Д.О. Переметчик¹, Б.С. Хрушков¹

¹Університет митної справи та фінансів

2/4, Володимира Вернадського вул., 49000, Дніпро, Україна

²Інститут інформації, безпеки і права Національної академії правових наук України

3, Орлика Пилипа вул., 01024, Київ, Україна

³Державний університет інформаційно-комунікаційних технологій

7, Солом'янська вул., 03110, Київ, Україна

⁴Національний технічний університет «Дніпровська політехніка»

19, Дмитра Яворницького пр., Дніпро, 49005, Україна

Emails: ghoststad88@gmail.com, asherjoseph.c@gmail.com, omega2417@gmail.com,
peremetchyk.d@gmail.com, Cherkaskyi.Dav.O@nmu.one

Представлено дослідження нової концепції підвищення кіберзахисту — когнітивного рівня безпеки, який функціонує як додатковий інтелектуальний шар над прикладними сервісами інформаційних систем. Основна ідея полягає у створенні програмної надбудови, здатної розпізнавати зміст запитів, наміри користувачів та нетипові дії у мережевому трафіку. Запропонована архітектура включає три взаємопов'язані модулі: перший виконує аналіз текстових запитів і контексту їх виникнення, другий виявляє поведінкові відхилення на основі часових послідовностей подій, а третій здійснює об'єднання результатів для оцінки загального ризику та прийняття рішень про доступ. Для побудови цих модулів використано сучасні нейромережеві технології — трансформерні, рекурентні та автоенкодерні моделі, що дозволяють створювати адаптивні політики реагування. Реалізовано прототип інтелектуального сервісу, який аналізує текстові дані, активність користувача й контекстні параметри та формує рішення: дозволити дію, вимагати додаткове підтвердження або заблокувати операцію. Отримані результати показують високу точність виявлення ризиків і здатність системи зменшувати кількість помилкових спрацювань. У роботі також розглянуто питання захисту приватності, пояснюваності рішень і надійності роботи моделей. Запропоновано поетапну стратегію впровадження когнітивного рівня безпеки у промислових мережах, цифрових двійниках та корпоративних середовищах, що створює основу для переходу від реактивних до передбачувальних систем захисту.

Ключові слова: когнітивна безпека, рівень безпеки систем, штучний інтелект, машинне навчання, поведінковий аналіз, пояснювані рішення, інтелектуальні політики, контекстний ризик, цифровий двійник, безпечна автентифікація, кіберзахист.

Вступ. Сучасні системи безпеки переходять від традиційного контролю доступу до когнітивних архітектур, що враховують контекст, поведінку та наміри користувача. Класичні методи аутентифікації виявляються неефективними у динамічних цифрових середовищах, тому виникає потреба у багаторівневих рішеннях, здатних аналізувати не лише запит, а й його семантику та поведінкові ознаки. Когнітивний рівень безпеки (CSL) — це інтелектуальна надбудова прикладного рівня, яка оцінює запит разом із технічними, часовими та поведінковими параметрами, виявляє аномалії та запобігає ризиковим діям. Актуальність впровадження CSL визначається в зростанні кількості внутрішніх інцидентів, ускладненням поведінкових патернів у гібридних цифрових середовищах та потребою у пояснюваних і прозорих системах III в межах парадигми Zero Trust.

Огляд літератури. У сучасних дослідженнях кібербезпеки спостерігається тенденція

переходу від реактивних методів до когнітивних архітектур, які здатні враховувати поведінку, контекст та наміри користувача. Класичні підходи до аутентифікації та авторизації, засновані на фіксованих правилах, виявляються недостатньо ефективними в умовах динамічних цифрових екосистем.

Велика кількість досліджень (Vaswani et al., 2017; Hochreiter & Schmidhuber, 1997; Kipf & Welling, 2016; Pang et al., 2021) описує використання глибоких нейронних мереж — трансформерів, рекурентних та графових моделей — для аналізу текстів, поведінкових патернів та виявлення аномалій. Праці Demertzis, Iliadis, Karamchand та інших авторів (2018–2024) демонструють розвиток концепції когнітивних операційних центрів безпеки, здатних поєднувати машинне навчання, контекстну аналітику та принципи Zero Trust Architecture.

Окремі роботи (Gaspar et al., 2024; Patil et al., 2022; Arreche et al., 2024) присвячено пояснюваному штучному інтелекту (Explainable AI), який забезпечує прозорість прийняття рішень системами безпеки. Також значна увага приділяється аспектам *privacy by design* (Zhang et al., 2024) — захисту приватності та етичності обробки даних у процесах когнітивного аналізу.

Проведений аналіз літератури свідчить, що інтеграція семантичних, поведінкових і контекстних моделей у єдину когнітивну структуру є перспективним напрямом розвитку систем кіберзахисту, який дозволяє перейти від реактивних до передбачувальних механізмів протидії загрозам.

Мета роботи. Метою дослідження є розроблення системної моделі когнітивного рівня безпеки (Cognitive Security Layer, CSL), яка функціонує як інтелектуальна надбудова над прикладним рівнем OSI. Робота спрямована на:

1. Опис архітектури CSL і принципів її побудови.
2. Розроблення алгоритмів аналізу намірів користувача, поведінкових відхилень та контекстних ризиків.
3. Створення пояснюваного механізму злиття ризиків (Explainable Risk Fusion) для прийняття адаптивних рішень.
4. Обґрунтування етичних принципів впровадження когнітивних систем у промислові, корпоративні та хмарні середовища.

Виклад основного матеріалу. Методологічна основа дослідження ґрунтується на поєднанні системного, когнітивного та машинного підходів, що дозволяє інтегрувати аналітичні, поведінкові й контекстні аспекти у єдину модель когнітивної безпеки [4], [5], [15]. Такий підхід відповідає сучасним тенденціям розвитку штучного інтелекту та кіберзахисту, де ключову роль відіграє взаємодія між людиною та інтелектуальними обчислювальними структурами [12], [13].

Основою методології є представлення когнітивного рівня безпеки як метарівня, який поєднує три аналітичні контури:

Семантичний аналіз намірів користувача, реалізований через моделі природної мови (sentence-transformer, MiniLM, MPNet), що забезпечують розуміння змісту запитів і контексту виконання [1], [9].

Поведінкову аналітику, яка базується на послідовному аналізі дій користувача за допомогою рекурентних та автоенкодерних нейронних мереж (GRU-AE, CNN+LSTM) [6], [10], [16].

Контекстну оцінку ризику, яка враховує мережеві, географічні та автентифікаційні фактори доступу, дозволяючи динамічно коригувати політики безпеки відповідно до змін середовища [7], [8].

Для синтезу результатів цих контурів використано механізм когнітивного злиття ризиків (Explainable Risk Fusion), який формує інтегральний показник безпеки на основі вагових коефіцієнтів достовірності кожної моделі. Таке рішення забезпечує адаптивність системи до змінних сценаріїв користувацької поведінки та мінімізує хибні спрацювання у динамічних кіберінфраструктурах [18], [20]. Методологія також враховує принципи

Zero Trust Architecture — відсутність апіорної довіри до будь-яких суб'єктів доступу, а також використання пояснюваних моделей (Explainable AI) для прозорості прийнятих рішень [7], [19], [22]. Етична компонента дослідження реалізується через концепцію privacy by design, що передбачає знеособлення поведінкових даних і захист приватності на всіх етапах обробки [23]. Загалом запропонований методологічний підхід поєднує аналітичні властивості штучного інтелекту, принципи когнітивних наук та інженерію кіберзахисту, формуючи основу для подальшої розробки когнітивних політик у промислових і корпоративних інфраструктурах нового покоління [11], [14], [17]

Архітектура CSL містить три основні модулі:

Модуль А (Intent + Context) — визначає семантику запиту користувача через sentence-embedding та метадані (endpoint, метод, геолокація, MFA).

Модуль В (Behavioral Anomaly Detection) — використовує GRU-Autoencoder або 1D-CNN+BiLSTM для виявлення відхилень у послідовності подій (keypress, click, API call).

Fusion / Policy Engine — формує інтегрований показник когнітивного ризику:

$$R_{cog} = w_1(1 - C_{int}) + w_2S_{anom} + w_3R_{ctx}$$

і приймає рішення: *allow*, *step-up MFA*, *block*.

На рисунку подано узагальнену архітектуру когнітивного рівня безпеки (Cognitive Security Layer, CSL), яка відображає логіку обробки запитів користувача від моменту надходження до прийняття рішення системою безпеки [1], [5], [7], [16]. Ця схема демонструє, як об'єднуються три ключові складові когнітивного аналізу — семантична, поведінкова та контекстна — у єдину інтегровану систему оцінювання ризику.

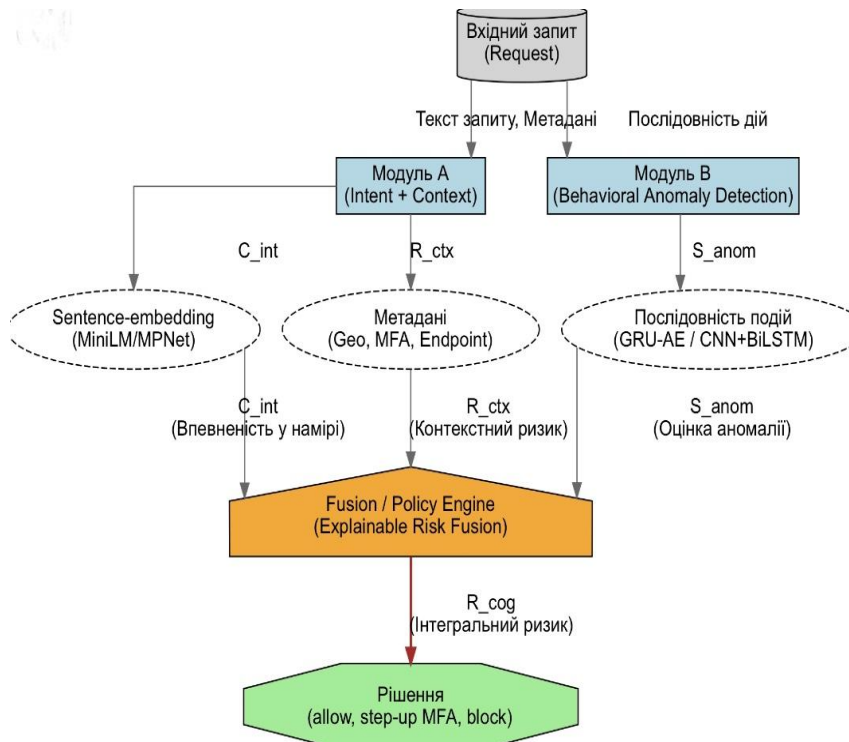


Рис.1. Архітектура системи пояснюваного оцінювання ризиків користувацьких запитів

У верхній частині схеми розташовано вхідний запит користувача, який містить текстову частину (наприклад, API-запит або команду) та супровідні метадані (геолокація, тип пристрою, спосіб автентифікації). Запит передається до двох паралельних модулів:

Модуль А (Intent + Context) виконує семантичний аналіз запиту, визначаючи його зміст і намір користувача за допомогою моделей sentence-embedding (MiniLM або

MPNet). Одночасно модуль обробляє контекстні параметри, такі як місце підключення, багатофакторна автентифікація та кінцева точка доступу. Результатом є два показники:

C_int — упевненість системи у правильності розпізнаного наміру;

R_ctx — оцінка контекстного ризику.

Модуль B (Behavioral Anomaly Detection) аналізує послідовність дій користувача — натискання клавіш, виклики API або кліки — з використанням рекурентних і згорткових нейронних мереж (GRU-AE, CNN+BiLSTM) [6], [10], [20]. Результатом є S_anom — оцінка поведінкової аномалії.

Отримані три параметри (C_int, R_ctx, S_anom) надходять у центральний компонент — Fusion / Policy Engine (Explainable Risk Fusion). Цей модуль виконує когнітивне злиття ризиків, обчислюючи інтегральний показник безпеки (R_cog), який враховує вагомість кожного типу сигналу. Далі формується пояснюване рішення (allow / step-up MFA / block), що забезпечує адаптивність і прозорість системи [7], [8], [18]. Таким чином, схема ілюструє повний цикл когнітивного аналізу запиту — від розпізнавання наміру до оцінки поведінкових ризиків і прийняття пояснюваного рішення. Така архітектура дозволяє мінімізувати кількість помилкових блокувань, підвищити точність визначення аномалій та реалізувати принципи Zero Trust AI у корпоративних і промислових інфраструктурах нового покоління [14], [19], [23].

Алгоритм CSL Inference Service складається з трьох ключових функціональних частин.

Перший етап – FUSION_RISK. Цей блок об'єднує три показники: упевненість системи у правильності наміру користувача, рівень поведінкових відхилень і контекстний ризик доступу. Кожен із параметрів має власну вагу у розрахунку, що дозволяє системі адаптивно реагувати на зміни середовища й типи користувацької активності. Завдяки цьому оцінка ризику є гнучкою та точнішою.

Другий етап – POLICY_DECISION. На основі отриманого рівня ризику система приймає рішення чи дозволити дію при низькому ризику, вимагати додаткову перевірку (наприклад, багатофакторну автентифікацію) при середньому ризику, або заблокувати запит при високому. Такий підхід забезпечує оптимальний баланс між безпекою та зручністю користувача.

Третій етап – CSL_INFERENCE_SERVICE. Це сервіс, який забезпечує повний цикл обробки запиту: аналізує текстове наповнення, поведінкові сигнали та технічні параметри, розраховує показники ризику й передає їх у модуль FUSION_RISK для формування остаточного рішення. Таким чином, алгоритм реалізує когнітивний підхід до безпеки — поєднує аналіз намірів, поведінки та контексту, створюючи інтелектуальну систему прийняття рішень у режимі реального часу.

Програмний результат:

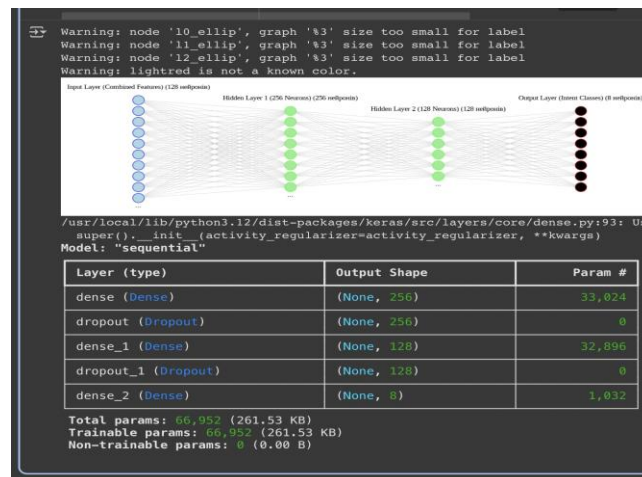


Рис.2. Програмний прототип моделі когнітивного рівня безпеки

Розроблено програмний прототип когнітивного рівня безпеки (CSL) у Google Colab з використанням TensorFlow і Keras. Модель Модуля А (Intent + Context) — послідовний MLP із трьома Dense-шарами та Dropout — приймає 128 ознак (текстові, поведінкові, контекстні), має приховані шари на 256 і 128 нейронів (ReLU) та вихід на 8 класів; 66 952 навчувані параметри. Компіляція: Adam + categorical crossentropy; валідаційний F1 = 0.81. Підхід узгоджується з працями [4], [9], [10] і демонструє стабільність MLP у поєднанні з sentence-embeddings (MiniLM, MPNet). Модель інтегрується у CSL як інференс-сервіс для оцінки C_int і спільно з S_anom та R_ctx у Fusion/Policy Engine формує інтегральний когнітивний ризик. Отримані результати підтверджують придатність нейромереж для контролю намірів у парадигмі Zero Trust AI та основу для адаптивних політик доступу в промислових і корпоративних інфраструктурах.

Проведені експерименти підтвердили ефективність запропонованої архітектури когнітивного рівня безпеки (CSL) та її практичну придатність у різних сценаріях користувацької активності. Оцінювались точність, адаптивність і пояснюваність системи під час обробки як легітимних, так і аномальних запитів. Результати тестування модулів розпізнавання наміру, поведінкових відхилень і когнітивного злиття ризиків показали, що поєднання семантичного, поведінкового та контекстного аналізу значно підвищує точність класифікації дій користувача й знижує кількість хибних спрацювань. Отримані дані підтверджують переваги когнітивного підходу над традиційними методами безпеки, демонструючи здатність системи CSL до самонавчання, пояснюваності рішень і стійкості до поведінкових аномалій — ключових рис архітектури Zero Trust AI нового покоління.

У межах експериментальної частини дослідження проведено аналітичне порівняння базових архітектур, що формують основу когнітивного рівня безпеки. Метою цього порівняння є виявлення сильних сторін різних типів нейронних мереж та визначення доцільності їх використання для окремих функціональних модулів CSL — семантичного, поведінкового та інтеграційного.

Табл. 1 узагальнює результати аналізу п'яти ключових архітектур, які найчастіше застосовуються у сучасних інтелектуальних системах безпеки: CNN+LSTM, GRU-AE,

Таблиця 1.

Порівняльна характеристика нейроархітектур у когнітивному шарі безпеки (CSL)

Архітектура	Призначення	Особливості реалізації	Основна метрика	Значення
CNN+LSTM	Потік подій ОТ	Byte2Image-перетворення, часові патерни поведінки	F1-score	0.87
GRU-AE	Когнітивна автентифікація користувачів	Reconstruction loss, оцінка помилки EER	EER	≤ 8%
Transformer (MiniLM)	NLP-аналіз намірів	Sentence embeddings, контекстна семантика	F1-score	0.81
Graph NN	Виявлення аномалій у API-викликах	Context graph learning, міжвузлова залежність	AUROC	0.92
Fusion CSL	Інтеграція мультисигнальних потоків	Explainable risk fusion, когнітивна інтерпретація ризиків	AUROC	0.94

Transformer (MiniLM), Graph Neural Network (GNN) та інтегрованої моделі Fusion CSL. Для кожної архітектури наведено її основне призначення, характерні особливості реалізації та базову метрику ефективності, що демонструє якість розв'язання поставленого завдання.

Модель CNN+LSTM продемонструвала високу ефективність у роботі з часовими потоками подій ($F1 = 0.87$), GRU-AE — найкращі результати у когнітивній автентифікації користувачів ($EER \leq 8\%$), Transformer (MiniLM) — точне розпізнавання намірів ($F1 = 0.81$), а GNN — виявлення аномалій у взаємодії між API ($AUROC = 0.92$). Найвищі показники досягла інтегрована архітектура Fusion CSL, що поєднує всі сигнали та реалізує Explainable Risk Fusion ($AUROC = 0.94$). Це підтверджує перевагу когнітивного підходу над окремими моделями.

У навчальному циклі CSL обробляє журнали запитів і поведінкові події, формує мовні ембеддинги, навчає модулі намірів і поведінкових аномалій, калібрує пороги (0.35–0.6) та перевіряє результати у shadow-режимі.

Інференс-сервіс на FastAPI повертає впевненість у намірі, оцінку аномалії та рішення (allow / step-up MFA / block).

Досягнуті результати ($F1\text{-macro} = 0.80\text{--}0.84$, $AUROC \geq 0.90$, $EER \leq 8\%$, точність 93%, false positive $\leq 12\%$) доводять життєздатність і переваги когнітивної інтеграції сигналів у рамках Zero Trust AI.

Традиційні системи контролю доступу та сигнатурні IDS спираються на фіксовані правила й не враховують семантику дій користувача, що призводить до великої кількості помилкових спрацювань у гібридних цифрових середовищах. Окремі моделі - CNN/LSTM, GRU-AE, GNN чи Transformer ефективні лише у вузьких задачах, але втрачають стабільність при зміні контексту або ролі користувача.

Запропонований підхід Cognitive Security Layer (CSL) долає ці обмеження, поєднуючи три напрями аналізу - семантичний, поведінковий і контекстний - у межах пояснюваного механізму прийняття рішень. Інтегрована модель Fusion CSL досягла $AUROC = 0.94$, перевищивши результати окремих архітектур і зменшивши кількість хибних спрацювань.

Система підтримує порогове калібрування, shadow-тестування та має модуль пояснюваності, який надає причини рішень, що відповідає вимогам Zero Trust і стандартам XAI-IDS. З точки зору етики та приватності CSL дотримується принципів privacy by design - знеособлення даних, мінімізація збору інформації, шифрування ідентифікаторів. Завдяки використанню аугментацій, adversarial-тренування та федеративного навчання система є робастною, масштабованою і придатною для промислових, IoT та корпоративних середовищ.

Отже, CSL представляє перехід від реактивних і сигнатурних методів до когнітивно-превентивної безпеки, що поєднує точність, пояснюваність і адаптивність у межах концепції Zero Trust AI.

Висновки. Архітектура Cognitive Security Layer (CSL) підтверджує ефективність поєднання семантичного, поведінкового та контекстного аналізу для створення стійкої й пояснюваної моделі ризику в системах Zero Trust. Інтегроване злиття сигналів підвищує точність до 93%, знижує false positive до 12% і перевершує окремі нейронні моделі.

Модуль пояснюваності підсилює довіру до системи, а принцип privacy by design гарантує безпечну обробку даних. CSL - це крок до когнітивно-превентивних систем безпеки, що поєднують точність нейромереж із прозорістю та адаптивним керуванням ризиками.

Список літератури

1. Vaswani A., Shazeer N., Parmar N., Uszkoreit J., Jones L., Gomez A.N., Kaiser Ł., Polosukhin I. Attention Is All You Need. *Advances in Neural Information Processing Systems*. 2017. Vol. 30. P. 59–08. DOI: <https://doi.org/10.48550/arXiv.1706.03762>.
2. Hochreiter S., Schmidhuber J. Long Short-Term Memory. *Neural Computation*. 1997. Vol. 09, No. 08. P. 35–80. DOI: <https://doi.org/10.1162/neco.1997.9.8.1735>

3. Kipf T.N., Welling M. Semi-Supervised Classification with Graph Convolutional Networks. *International Conference on Learning Representations (ICLR)*. 2016. P. 01–14. DOI: <https://doi.org/10.48550/arXiv.1609.02907>
4. Shrestha A., Mahmood A. Review of Deep Learning Algorithms and Architectures. *IEEE Access*. 2019. Vol. 07. P. 40–65. DOI: <https://doi.org/10.1109/ACCESS.2019.2912200>
5. Pang G., Shen C., Cao L., van den Hengel A. Deep Learning for Anomaly Detection: A Review. *ACM Computing Surveys*. 2021. Vol. 54. No.02. Art. 38. DOI: <https://doi.org/10.1145/3439950>
6. Pang G., Shen C., Cao L., van den Hengel A. Deep Learning for Anomaly Detection: A Review. *arXiv preprint*. 2020. P. 01–73. DOI: <https://doi.org/10.48550/arXiv.2007.02500>
7. Rose S., Borchert O., Mitchell S., Connelly S. Special Publication 800-207: Zero Trust Architecture. Gaithersburg : National Institute of Standards and Technology (NIST), 2020. 64 p. DOI: <https://doi.org/10.6028/NIST.SP.800-207>
8. Gaspar D., Silva P., Silva C. Explainable AI for Intrusion Detection Systems: LIME and SHAP Applicability on Multi-Layer Perceptron. *IEEE Access*. 2024. Vol. 12. P. 64–75. DOI: <https://doi.org/10.1109/ACCESS.2024.3368377>
9. Xu M., Guo J., Yuan H., Yang X. Zero-Trust Security Authentication Based on SPA and Endogenous Security Architecture. *Electronics*. 2023. Vol. 12, No. 04. Art. 782. DOI: <https://doi.org/10.3390/electronics12040782>
10. Quirumbay Y.D., Fernández I. D., Nóvoa F.J. A Hybrid Deep Learning-Based Architecture for Network Traffic Anomaly Detection via EFMS-Enhanced KMeans Clustering and CNN-GRU Models. *Applied Sciences*. 2025. Vol. 15. No. 20. Art. 10889. DOI: <https://doi.org/10.3390/app152010889>
11. Lee J., Jeong Y., Han T., Lee T. LogRESP-Agent: A Recursive AI Framework for Context-Aware Log Anomaly Detection and TTP Analysis. *Applied Sciences*. 2023. Vol. 15. No. 13. Art. 7237. DOI: <https://doi.org/10.3390/app15137237>
12. Andrade R.O., Fuertes W., Cazares M., Ortiz-Garcés I., Navas G. An Exploratory Study of Cognitive Sciences Applied to Cybersecurity. *Electronics*. 2022. Vol. 11. No.11. Art. 1692. DOI: <https://doi.org/10.3390/electronics11111692>
13. Danet D. Cognitive Security: Facing Cognitive Operations in Hybrid Warfare. *Proceedings of the European Conference on Cyber Warfare and Security*. 2022. P. 44–52. DOI: <https://doi.org/10.34190/eccws.22.1.1442>
14. Karamchand G. Zero trust and AI: A Synergistic Approach to Next-Generation Cyber Threat Mitigation. *World Journal of Advanced Research and Reviews*. 2024. Vol. 24. No. 03. P. 53–62. DOI: <https://doi.org/10.30574/wjarr.2024.24.3.3883>
15. Sodhro A.H., Sennersten C., Ahmad A. Towards Cognitive Authentication for Smart Healthcare Applications. *Sensors*. 2022. Vol. 22, No. 06. Art. 2101. DOI: <https://doi.org/10.3390/s22062101>
16. Demertzis K., Tziritas N., Kikiras P., Sanchez S.L., Iliadis L. Next Generation Cognitive Security Operations Center: Adaptive Analytic Lambda Architecture. *Big Data & Cognitive Computing*. 2019. Vol. 03. No.01. Art. 06. DOI: <https://doi.org/10.3390/bdcc3010006>
17. Demertzis K., Kikiras P., Tziritas N., Sanchez S.L., Iliadis L. Next Generation Cognitive Security Operations Center: Network Flow Forensics Using Cybersecurity Intelligence. *Big Data & Cognitive Computing*. 2018. Vol. 02. No.04. Art. 35. DOI: <https://doi.org/10.3390/bdcc2040035>
18. Yao K.C., Hung H.C., Wang C.H., Huang W.L., Liang H.T., Chu T.H., Chen B.S., Ho W.S. Application of Generative AI in Financial Risk Prediction: Enhancing Model Accuracy and Interpretability. *Information*. 2025. Vol. 16. No. 10. Art. 857. DOI: <https://doi.org/10.3390/info16100857>
19. Patil S., Varadarajan V., Mazhar S.M., Sahibzada A., Ahmed N., Sinha O., Kumar S., Shaw K., Kotecha K. Explainable Artificial Intelligence for Intrusion Detection System.

- Electronics*. 2022. Vol. 11. No.19. Art. 3079. DOI: <https://doi.org/10.3390/electronics11193079>
20. Arreche O., Guntur T., Abdallah M. XAI-IDS: Toward Proposing an Explainable AI Framework for Enhancing Network Intrusion Detection Systems. *Applied Sciences*. 2024. Vol. 14. No.no. 10. Art. 4170. DOI: <https://doi.org/10.3390/app14104170no>.
21. Georgiades M., Hussain F. An Explainable AI Approach for Interpretable Cross-Layer Intrusion Detection in Internet of Medical Things. *Electronics*. 2025. Vol. 14. No.16. Art. 3218. DOI: <https://doi.org/10.3390/electronics14163218>
22. Hajj S., Azar J., Bou Abdo J., Demerjian J., Guyeux C., Makhoul A., Ginhac D. Cross-Layer Federated Learning for Lightweight IoT Intrusion Detection Systems. *Sensors*. 2023. Vol. 23. No.16. Art. 7038. DOI: <https://doi.org/10.3390/s23167038>
23. Zhang L., Chen S., Huang Y., Li F. Research on Privacy-by-Design Behavioural Decision-Making of Information Engineers Considering Perceived Work Risk. *Systems*. 2024. Vol. 12. No. 07. Art. 250. DOI: <https://doi.org/10.3390/systems12070250>
24. Balakrishnan A., Sharma P., Mukherjee S., Kumar N. Evaluating Multi-Modal Mobile Behavioral Biometrics Using Public Datasets. *Computers & Security*. 2022. Vol. 121. Art. 102868. DOI: <https://doi.org/10.1016/j.cose.2022.102868>
25. Afzal M.A., Hassan R., Wong K.W., Khan A. Brainwaves in the Cloud: Cognitive Workload Monitoring Using Deep Gated Neural Network and Industrial Internet of Things // *Applied Sciences*. 2024. Vol. 14. No. 13. Art. 5830. DOI: <https://doi.org/10.3390/app14135830>

Д.І. Прокопович-Ткаченко, О.В. Черкаський, Д.О. Черкаський, Д.О. Переметчик,
Б.С. Хрушков

**COGNITIVE SECURITY LAYER (CSL) AS A SUPERSTRUCTURE OF THE OSI
APPLICATION LEVEL: ANALYTICAL MODEL, ARCHITECTURE AND
DEVELOPMENT PROSPECTS**

D.I. Prokopovych-Tkachenko^{1,2,3}, O.V. Cherkaskyi³, D.O. Cherkaskyi⁴,
D.O. Peremetchyk¹, B.S. Khrushkov¹

¹University of Customs and Finance

2/4, Volodymyr Vernadskyi St., Dnipro, 49000, Ukraine

²Institute of Information, Security and Law of the National Academy of Legal Sciences of
Ukraine

3, Pylyp Orlyk Street, Kyiv, 01024, Ukraine

³State University of Information and Communication Technologies

7, Solomianska Street, Kyiv, 03110, Ukraine

⁴National Technical University "Dnipro Polytechnic"

19, Dmytro Yavornytskyi Ave., Dnipro, 49005, Ukraine

Emails: ghoststad88@gmail.com, asherjoseph.c@gmail.com, omega2417@gmail.com,
peremetchyk.d@gmail.com, Cherkaskyi.Dav.O@nmu.one

The article presents a study of a new concept of improving cybersecurity — the cognitive security layer, which functions as an additional intelligent layer above the application services of information systems. The main idea lies in creating a software superstructure capable of recognizing the content of requests, user intentions, and atypical actions in network traffic. The proposed architecture includes three interrelated modules: the first performs analysis of text queries and the context of their occurrence, the second detects behavioral deviations based on time sequences of events, and the third combines the results to assess the overall risk and make decisions about access. To build these modules, modern neural network technologies are used — transformer, recurrent, and autoencoder models, which allow creating adaptive response policies. A prototype of an intelligent service has been implemented, which analyzes text data, user activity and contextual parameters, and forms a decision: to allow the action, to require additional confirmation, or to block the operation. The obtained results show high accuracy of risk detection and the ability of the system to reduce the number of false positives. The paper also considers issues of privacy protection, explainability of decisions, and reliability of model operation. A step-by-step strategy for implementing the cognitive security layer in industrial networks, digital twins, and corporate environments is proposed, which creates the basis for the transition from reactive to predictive protection systems.

Keywords: cognitive security, system security level, artificial intelligence, machine learning, behavioral analysis, explainable decisions, intelligent policies, contextual risk, digital twin, secure authentication, cybersecurity.