

блокувати виконання програми. Узагальнену схему взаємодії класу з Telegram API, сервісами SerpAPI і imgbb та файловою системою подано на рисунку 1 та наведено основні потоки даних і місця збереження результатів.

На основі проведеного дослідження розроблено архітектуру системи верифікації відкритих джерел, яка охоплює декілька взаємопов'язаних модулів. Усі ці елементи об'єднано у єдину систему інтегрованої оцінки, що забезпечує можливість об'єктивного визначення ступеня достовірності досліджуваного матеріалу. Важливою складовою системи є застосування Source Credibility Matrix, яка дозволяє кількісно характеризувати надійність каналу на основі його історії, активності, стабільності публікацій та виявлених ознак маніпулятивності. Дані елементи створюють теоретично обґрунтований інструмент для оцінювання джерел інформації, що особливо важливо в умовах інформаційних загроз, спрямованих на українське суспільство.

1. What is OSINT: Open-source intelligence? European data. URL: <https://data.europa.eu/en/publications/datastories/what-osint-open-source-intelligence>.
2. Як ефективно використовувати OSINT-інструментарій? Команда «OsintFlow» ділиться досвідом // Державна прикордонна служба України. URL: <https://dpsu.gov.ua/uk/news/43174-YAk-efektivno-vikoristovuvati-OSINT-instrumentariy-Komanda-OsintFlow-dilitsya-dosvidom>.

### **Сучасні підходи до трансформації систем охорони праці на основі штучного інтелекту та предикативної аналітики**

УДК 004.056.5:57.087.1 Михайло Пригара<sup>1</sup>, В'ячеслав Шматуха<sup>2</sup>, Володимир Щербина<sup>3</sup>

<sup>1</sup>Ужгородський національний університет, [misha\\_prigara@ukr.net](mailto:misha_prigara@ukr.net),

<sup>2</sup>Київська школа економіки, [vvshmatukha@gmail.com](mailto:vvshmatukha@gmail.com),

<sup>3</sup>Державний університет «Київський авіаційний інститут», [smya@kai.edu.ua](mailto:smya@kai.edu.ua)

Традиційні системи управління охороною праці (СУОП) тривалий час базувалися на реактивному підході — аналізі інцидентів, що вже відбулися. Однак у 2025–2026 роках девіз Всесвітнього дня охорони праці «Революційні підходи до здоров'я і безпеки: роль ШІ та цифровізації» підкреслив глобальний перехід до проактивних стратегій. Використання ШІ дозволяє не лише фіксувати порушення, а й передбачати їх, створюючи динамічне безпечне середовище.

#### **1. Ключові напрями використання ШІ в ОП**

Сучасні підходи можна класифікувати за технологічними доменами:

- Комп'ютерний зір (Computer Vision) для моніторингу в реальному часі
- Це найбільш поширений напрям у 2026 році. Камери з ШІ інтегруються в існуючі системи відеоспостереження для:
- Автоматичного контролю використання засобів індивідуального

захисту (ЗІЗ): касок, жилетів, масок, рукавичок.

- Детекції небезпечних зон: сповіщення працівника та оператора при вході людини в зону роботи кранів або навантажувачів.
- Аналізу ергономіки: відстеження рухів працівника для запобігання скелетно-м'язовим розладам через неправильні пози чи перевантаження.

#### Б. Предиктивна аналітика та прогнозування ризиків

Алгоритми машинного навчання обробляють великі масиви даних (Big Data), включаючи звіти про дрібні інциденти (near-misses), погодні умови, стан обладнання та навіть психофізіологічні показники працівників. ШІ виявляє патерни, що передують аваріям, дозволяючи менеджменту втрутитися до моменту виникнення нещасного випадку.

#### В. Інтеграція з носимими пристроями (Wearables) та IoT

Розумні годинники та сенсори на одязі моніторять життєво важливі показники (пульс, температура тіла, рівень втоми). Це критично для робіт у замкнених просторах, при екстремальних температурах або у нічні зміни. При виявленні критичного рівня втоми система рекомендує зробити перерву.

#### 2. Трансформація навчання та інструктажів

ШІ змінює підхід до підготовки персоналу:

- Адаптивне навчання: Генеративний ШІ створює персоналізовані курси на основі помилок, які працівник допускав раніше.
- VR/AR-симуляції: Тренування навичок безпечної роботи у віртуальному середовищі з ШІ-інструктором, який моделює критичні ситуації в режимі реального часу.

#### 3. Переваги та економічна ефективність

За даними міжнародних звітів 2025-2026 рр., підприємства, що впровадили ШІ-платформи безпеки, демонструють:

- Зниження травматизму на 25–40%.
- Скорочення витрат на страхові виплати та компенсації.
- Прискорення аудитів: підготовка до перевірок з охорони праці стає швидшою на 40% завдяки автоматизованому збору даних.

#### 4. Виклики та етичні аспекти

Незважаючи на технологічний прогрес, залишаються відкритими питання:

- Конфіденційність даних: необхідність балансу між моніторингом безпеки та приватністю працівника.
- Психологічний тиск: ризик виникнення стресу у персоналу через відчуття «постійного нагляду» алгоритмом.
- Технологічна залежність: ризик втрати навичок самостійної оцінки безпеки людиною.

Висновок. Станом на 2026 рік штучний інтелект перестав бути футуристичним концептом і став фундаментальним інструментом у сфері охорони праці. Перехід від «контролю після факту» до «запобігання до події» є ключовим вектором розвитку. Майбутнє СУОП полягає у синергії людського досвіду та обчислювальної потужності ШІ, де технології виступають не як заміна інженеру з ОП, а як його високотехнологічний асистент.

1. Healthy Workplaces Summit 2025: discover key takeaways, photos and resources on safe digital work / EU-OSHA. Bilbao, 2025. URL: <https://osha.europa.eu/en/highlights/healthy-workplaces-summit-2025-discover-key-takeaways-photos-and-resources-safe-digital-work> (дата звернення: 06.05.2026).
2. Mishiba, Takenori. 2024. "Transforming Occupational Health and Safety Regulation: Strategic Pathways in the Era of Industry 4.0." *Journal of Occupational Health Law and Emerging Vision* 3, no. 2: 151–169. <https://doi.org/10.57523/jaohlev.pp.24-016> (дата звернення: 06.05.2026).
3. World Day for Safety and Health at Work 2025: Revolutionizing Health and Safety: The Role of AI and Digitalization at Work / International Labour Organization (ILO). 2025. URL: <https://www.ilo.org/safeday> (дата звернення: 06.05.2026).

### **Mitigating AI-driven security risks in educational software systems**

UDK 621.395.7 (043.2)

Stepan Prokipchyn<sup>1</sup>

*State University of Information and Communication Technologies,  
s.prokipchyn@stud.duikt.edu.ua*

The growing use of autonomous AI agents in educational software introduces new cybersecurity challenges due to their ability to interact with external systems and act on behalf of users. The objective of this work is to analyze access control-related risks in such systems and propose practical mitigation strategies. The relevance of the study is driven by the increasing integration of AI into critical educational processes. The scientific novelty is the structured analysis of these risks across different layers of AI usage within educational systems.

The vulnerability surface of AI systems is often defined as prompt injection, data poisoning and hallucination [1, 2]. However, any LLM is prone to these kinds of risks. What makes agentic systems especially vulnerable to these attacks is the main strength of the ReAct pattern – the ability to interact with external systems (load data, perform actions).

In the context of educational systems, the work covers three classes of AI agent use. Internal automation agents – purpose-built agents that automates operational scenarios, running scripted flows (e.g., automated syllabus review per instructional design, suggesting additional materials for students based on their results, AI-assisted learning, etc.). AI-assisted coding and engineering – in 2026 this has become an industry standard, with a significant portion of code generated by AI agents. AI-generated code in system-critical domains like security can lead to vulnerabilities. AI features in the product – production-side features powered by AI agents that students and educators interact with directly. Exposing AI to actual users without proper guardrails not only allows the system to be tricked or abused by dishonest individuals but can also lead to unexpected costs and overall system instability. All three layers are prone to access control-related security risks which are not new but are rather elevated by AI [3].