

Застосування штучного інтелекту та нечіткої логіки для оцінювання правових ризиків кадрових рішень у системі кадрового менеджменту

Application of Artificial Intelligence and Fuzzy Logic for Assessing Legal Risks of Personnel Decisions in the Human Resource Management System

Юрій Гусак ^A

Corresponding author: доктор військових наук, професор, професор кафедри, e-mail: y.husak1512@gmail.com, ORCID ID: <https://orcid.org/0000-0002-3423-2112>

Денис Кандуєв ^B

кандидат юридичних наук, офіцер відділу, e-mail: kanduevdenis@gmail.com, ORCID ID: <https://orcid.org/0009-0000-5827-8599>

Yurii Husak ^A

Corresponding author: Doctor of Military Sciences, Professor, Professor of the Department, e-mail: y.husak1512@gmail.com, ORCID ID: <https://orcid.org/0000-0002-3423-2112>

Denis Kanduev ^B

Candidate of Legal Sciences, Officer of the Department, e-mail: kanduevdenis@gmail.com, ORCID ID: <https://orcid.org/0009-0000-5827-8599>

^A Національний університет оборони України, м. Київ, Україна

^B Головне управління супроводження ракетних програм Збройних Сил України, м. Київ, Україна

^A National Defence University of Ukraine, Kyiv, Ukraine

^B Main Directorate for Support of Missile Programs of the Armed Forces of Ukraine, Kyiv, Ukraine

Received: July 10, 2026 | Revised: August 21, 2026 | Accepted: August 31, 2026

УДК: 004.8:005.95/.96:355.1:340.13(477)

DOI: <https://doi.org/10.33445/sds.2026.16.4.15>

Мета дослідження. Розробити гібридну модель застосування штучного інтелекту та нечіткої логіки для оцінювання правових ризиків кадрових рішень у системі правового забезпечення кадрового менеджменту.

Методи дослідження. Використано системний аналіз, логіко-предикатне моделювання, ризик-орієнтований підхід, нечітку логіку, правила Mamdani, дефазифікацію методом центру ваги, технології обробки природної мови, Retrieval-Augmented Generation, машинне навчання, експертні системи, графи знань і концепцію Human-in-the-Loop.

Результати дослідження. Запропоновано архітектуру інтелектуальної системи оцінювання правових ризиків кадрових рішень, що поєднує модулі обробки кадрових даних, пошуку нормативних підстав, LLM/RAG-аналізу, експертну базу правил, граф знань, нечітку систему Mamdani, машинне навчання, засоби пояснюваності та людський контроль. Вхідними змінними визначено повноту документів, дотримання процедури, підтвердження повноважень, непорушення прав особи та захист персональних даних. На умовному прикладі показано зниження інтегрального правового ризику з $(R_{\{legal\}}=0,50)$ до $(R_{\{legal\}}=0,396)$ після покращення документального забезпечення та процедури.

Практична цінність дослідження. Модель може використовуватися для попереднього оцінювання кадрових рішень і формування пояснених рекомендацій за обов'язкового людського контролю.

Обмеження/майбутні дослідження. Модель має концептуально-методичний характер. Функції належності потребують валідації на реальних кадрових кейсах, а демонстраційна база правил Mamdani — розширення для різних типів кадрових рішень. Подальші дослідження передбачають створення програмного прототипу, формування датасету кадрово-правових кейсів та тестування точності AI-модулів.

Тип статті: методична.

Purpose. To develop a hybrid model combining artificial intelligence and fuzzy logic for assessing the legal risks of personnel decisions within the legal support system of human resource management.

Methods. The study employed systems analysis, logical-predicate modelling, a risk-oriented approach, fuzzy logic, Mamdani rules, centroid defuzzification, natural language processing technologies, Retrieval-Augmented Generation, machine learning, expert systems, knowledge graphs, and the Human-in-the-Loop concept.

Findings. An architecture of an intelligent system for assessing the legal risks of personnel decisions is proposed. It integrates modules for personnel data processing, retrieval of legal grounds, LLM/RAG-based analysis, an expert rule base, a knowledge graph, a Mamdani fuzzy inference system, machine learning, explainability tools, and human oversight. The input variables include document completeness, procedural compliance, confirmation of authority, protection of individual rights, and personal data protection. A conditional example demonstrates a decrease in the integral legal risk from $(R_{\{legal\}}=0.50)$ to $(R_{\{legal\}}=0.396)$ after improvements in documentary support and procedural compliance.

Practical implications. The model can be used for the preliminary assessment of personnel decisions and for generating explainable recommendations under mandatory human oversight.

Limitations/Future research. The model is conceptual and methodological in nature. The membership functions require validation using real personnel cases, while the demonstrative Mamdani rule base needs to be expanded for different types of personnel decisions. Future research should focus on developing a software prototype, creating a dataset of personnel-related legal cases, and testing the accuracy of AI modules.

Paper type: Methodological.

Ключові слова: штучний інтелект; нечітка логіка; кадрові рішення; правовий ризик; кадровий менеджмент; правове забезпечення; Mamdani; RAG; LLM; Human-in-the-Loop; персональні дані; Збройні Сили України.

Key words: Artificial Intelligence; Fuzzy Logic; Personnel Decision; Legal Risk; Human Resource Management; Legal Support; Mamdani; RAG; LLM; Human-in-the-Loop; Personal Data.

Вступ

Правове забезпечення кадрового менеджменту є необхідною умовою законності, обґрунтованості, прозорості та контрольованості кадрових рішень. У кадровій практиці значна частина рішень не може бути оцінена лише за бінарною логікою “правомірно / неправомірно”. Часто виникають проміжні стани: документи наявні, але не всі; процедура загалом дотримана, але має недоліки; правова підстава існує, але потребує уточнення; повноваження підтверджені, але є сумнів щодо процедури погодження; персональні дані обробляються, але доступ до них потребує додаткового обмеження.

У таких умовах правовий ризик кадрового рішення є не лише юридичною, а й інформаційно-аналітичною категорією. Він формується під впливом як формальних правових чинників, так і якості даних, повноти документів, надійності процедур, рівня захисту прав особи та стану інформаційної безпеки. Саме тому оцінювання правового ризику потребує поєднання правового аналізу, математичного моделювання, експертних правил і сучасних методів штучного інтелекту.

Штучний інтелект у цій сфері може виконувати не функцію автономного прийняття кадрового рішення, а функцію інтелектуальної підтримки: аналізувати документи, знаходити релевантні правові норми, виявляти неповні або суперечливі дані, класифікувати тип кадрової процедури, прогнозувати ймовірність оскарження, формувати пояснювані рекомендації та передавати результати для людської перевірки. Міжнародні підходи до регулювання ШІ також орієнтують на ризик-орієнтоване та людиноцентричне застосування інтелектуальних систем. Принципи OECD щодо ШІ спрямовані на розвиток інноваційного і надійного ШІ, який поважає права людини та демократичні цінності; вони були прийняті у 2019 році та оновлені у 2024 році.

Особливої уваги потребують системи ШІ, які застосовуються у сфері зайнятості, управління працівниками та кадрових рішень, оскільки вони можуть впливати на права, кар’єру, доступ до посад і професійний статус особи. У європейському регулюванні такі сфери належать до високоризикових або чутливих напрямів застосування ШІ.

Для сфери правового забезпечення кадрового менеджменту важливо, щоб ШІ не замінював юридичну експертизу, а працював у зв’язці з нечіткою експертною системою, правовою базою даних, графом знань і людиною-експертом. Нечітка логіка в такій системі є інструментом формалізації проміжних станів відповідності, а ШІ – інструментом вилучення ознак, аналізу текстів, пошуку норм і пояснення результатів.

У традиційних підходах оцінювання кадрового рішення часто зводиться до перевірки відповідності окремим вимогам: наявність документів, наявність правової підстави, повноваження посадової особи, дотримання процедури. Однак у реальній практиці ці умови не завжди мають бінарний характер. Наприклад, документи можуть бути подані частково; процедура може бути виконана із затримкою; персональні дані можуть оброблятися на законній підставі, але з недостатнім рівнем доступового контролю. Такі ситуації вимагають нечіткого, градуйованого оцінювання.

Нечітка логіка є одним із найбільш придатних інструментів для формалізації таких проміжних станів. Вона дозволяє визначати ступінь належності об’єкта до певної категорії, наприклад “низька відповідність”, “середня відповідність”, “висока відповідність”. У вихідному матеріалі нечітка логіка була застосована саме для таких ситуацій, коли правовий стан кадрового рішення не можна оцінити тільки як “законно” або “незаконно”.

Водночас сучасний розвиток ШІ дає змогу розширити класичну нечітку модель. Зокрема, програмні модулі ШІ для обробки природної мови (Natural Language Processing, NLP-

модулі) можуть автоматично вилучати з кадрових документів правові ознаки; модулі ШІ, які поєднують пошук інформації та генерацію відповіді (Retrieval-Augmented Generation, RAG-модулі) можуть знаходити релевантні нормативні положення; великі мовні моделі (Large Language Model, LLM) можуть формувати пояснювані юридичні довідки; машинне навчання може калібрувати ваги, пороги та функції належності на основі історичних кейсів; графи знань можуть пов'язувати особу, посаду, процедуру, документ, норму права та ризик.

Міжнародні документи також підкреслюють потребу в управлінні ризиками ШІ [1] – [5]. Так, рамковий документ Національного інституту стандартів і технологій США (NIST) AI Risk Management Framework [3] орієнтований на практичне управління ризиками ШІ-систем, а не лише на декларативні принципи.

Регламент Європейського Союзу про штучний інтелект (Regulation (EU) 2024/1689) (AI Act) [4] запроваджує ризик-орієнтовану логіку регулювання систем ШІ, а додаток III (Annex III) до AI Act [5] містить перелік високоризикових сфер застосування, серед яких згадуються напрями, пов'язані із зайнятістю та управлінням працівниками.

Отже, науково-практичне завдання полягає в побудові такої моделі, яка об'єднує можливості ШІ щодо обробки даних і текстів із перевагами нечіткої логіки щодо оцінювання невизначеності та часткової відповідності.

Метою статті є розроблення гібридної AI-Fuzzy моделі оцінювання правового ризику кадрових рішень для системи правового забезпечення кадрового менеджменту.

Для досягнення мети поставлено такі завдання:

визначити правовий ризик кадрового рішення як об'єкт інтелектуального та нечіткого моделювання;

сформувати архітектуру AI-Fuzzy системи оцінювання правових ризиків кадрових рішень;

визначити роль ШІ у вилученні ознак, пошуку нормативних підстав, аналізі документів і формуванні пояснень;

побудувати нечітку модель оцінювання правового ризику кадрового рішення;

сформувати базу правил Mamdani для оцінювання ризику;

провести приклад розрахунку для базового та покращеного сценаріїв;

визначити наукову новизну, практичну значущість, обмеження та напрями подальших досліджень.

Огляд літератури

Аналітичний огляд наукових публікацій за 2021–2026 роки з питань застосування ШІ у кадровому менеджменті свідчить, що у фахових виданнях ця проблематика розробляється переважно на якісному рівні та охоплює кілька основних напрямів: автоматизацію добору персоналу, оцінювання компетентностей, прогнозування плинності кадрів, планування кадрових потреб, аналіз продуктивності, цифровий профіль працівника, підготовку документів і підтримку управлінських рішень. Проте правовий ризик кадрового рішення як самостійний об'єкт математичного та ШІ-моделювання залишається недостатньо формалізованим.

Робота Shaout A. та Al-Shammari M. [6] є однією з ранніх робіт, у яких нечітка логіка застосовується до оцінювання персоналу. Автори показали, що оцінювання ефективності роботи персоналу містить значну частку суб'єктивних, лінгвістичних і неточних оцінок, тому нечітка логіка (Fuzzy logic) може бути корисною для формалізації таких показників, як “добрий рівень виконання”, “середня продуктивність”, “висока відповідальність” тощо. Ця публікація важлива тим, що вона демонструє придатність нечіткої логіки для HR-сфери. Проте його обмеження полягає в тому, що об'єктом оцінювання є працівник або його результативність, а не правовий ризик кадрового рішення. Тобто fuzzy-моделі в HR-моделюванні існують, але вони переважно спрямовані на оцінювання персоналу, а не на правове забезпечення кадрового менеджменту.

У роботі [7] Sun C. С. запропоновано модель оцінювання результативності на основі поєднання нечіткої логіки з методом аналізу ієрархій (Analytic Hierarchy Process, AHP) (метод Fuzzy AHP) і нечіткої логіки з методом вибору альтернативи (Technique for Order Preference by Similarity to Ideal Solution, TOPSIS) (метод Fuzzy TOPSIS). Fuzzy AHP використовується для визначення ваг критеріїв, а Fuzzy TOPSIS – для ранжування альтернатив. Sun C. С. працює з лінгвістичними оцінками, які параметризуються трикутними нечіткими числами, що дозволяє враховувати суб'єктивність, неточність і невизначеність експертних суджень. Ця публікація підтверджує придатність нечітких багатокритеріальних моделей для управлінського оцінювання. У Sun C. С. ідеться про оцінювання результативності, а не про оцінювання правового ризику кадрового рішення.

Робота Dursun M. та Karsak E. E. [8] належить до напряму застосування групи методів нечіткого багатокритеріального прийняття рішень (Fuzzy Multi-Criteria Decision Making, Fuzzy MCDM) (метод Fuzzy MCDM) у кадровому менеджменті. Автори розробили нечіткий багатокритеріальний алгоритм для добору персоналу, який використовує лінгвістичне представлення оцінок, агрегування нечіткої інформації та TOPSIS для вибору найкращої альтернативи. Публікація підтверджує те, що кадрові рішення можуть формалізуватися як багатокритеріальні задачі в умовах нечіткості. Разом з тим, предметом публікації є добір персоналу, а не правова перевірка кадрового рішення.

У цій статті [9] Kusumawardani R. P. та Agintiara M. методи Fuzzy AHP-TOPSIS застосовано для вибору менеджера з персоналу. Автори поєднали переваги AHP для визначення ваг критеріїв і TOPSIS для ранжування кандидатів, оскільки перевагою AHP є виведення ваг критеріїв через матрицю попарних порівнянь, а TOPSIS – здатність визначати найкращу альтернативу. В цій публікації показано як нечіткі методи можуть бути використані в HR-рішенні, де багато якісних критеріїв важко виміряти точно. Таким чином, нечіткі методи вже застосовуються в HR-менеджменті, але переважно для вибору або ранжування персоналу. Натомість застосування нечіткої логіки для оцінювання правового ризику кадрового рішення в органах військового управління залишається менш розробленим.

Найбільш близькою до теми статті є публікація Hoque S. та інших [10], яка поєднує великі мовні моделі, NLP та Fuzzy TOPSIS. Автори розробили систему автоматизованого добору персоналу, яка аналізує профілі кандидатів, використовує LLM/NLP для вилучення ознак, а Fuzzy TOPSIS – для багатокритеріального ранжування кандидатів. Ця праця демонструє сучасний тренд: поєднання генеративного або трансформерного ШІ з нечіткими методами прийняття рішень. Разом з тим, результатом застосування цих методів є рейтинг кандидатів, а не інтегральний правовий ризик кадрового рішення R_{legal}

Автори Varocas S. та Selbst A. D. [11] показують, що моделі великих даних можуть створювати дискримінаційний ефект навіть тоді, коли розробники не мають дискримінаційного наміру. Це особливо важливо для HR-менеджменту, оскільки кадрові рішення можуть спиратися на історичні дані, які вже містять упередження, що створює правові ризики застосування ШІ в кадровому менеджменті. В органах військового управління такі ризики можуть стосуватися нерівного доступу до посад, непрозорого оцінювання, необґрунтованого ранжування, упередженого прогнозування службової придатності або ризику оскарження кадрового рішення.

У документі Організації економічного співробітництва та розвитку (OECD) [12], який розроблений сформульовано міжнародні принципи відповідального та надійного використання штучного інтелекту (AI Principles). OECD визначає їх як п'ять ціннісних принципів і п'ять рекомендацій для політиків та учасників AI-екосистеми, які задають рамку, в якій ШІ має бути людиноцентричним, прозорим, підзвітним, надійним і таким, що поважає права людини. Тобто відповідати концепції Human-in-the-Loop [13], яка вимагає включення людини-експерта

до критичних етапів функціонування AI-системи, зокрема перевірки результатів, інтерпретації висновків і прийняття остаточного рішення. У військовому кадровому менеджменті ці принципи мають підвищене значення, оскільки кадрові рішення впливають на статус військовослужбовця, проходження служби, кар'єру, соціальні гарантії та персональні дані. А для органів військового управління це означає, що ШІ може використовуватися для аналізу кадрових документів, виявлення правових ризиків, пошуку нормативних підстав і підготовки рекомендацій, однак остаточне кадрове рішення повинно залишатися за уповноваженою посадовою особою.

У стратегіях НАТО щодо ШІ [14] та [15] започатковано системне бачення відповідального використання ШІ в обороні та закріплено шість принципів відповідального використання ШІ в обороні: правомірність (lawfulness), відповідальність і підзвітність (responsibility and accountability), пояснюваність і простежуваність (explainability and traceability), надійність (reliability), керованість і мінімізація упередженості (governability and bias mitigation). Тобто загальні AI-принципи перенесено у сферу військового управління. При цьому, такі принципи як правомірність, пояснюваність і простежуваність та відповідальність і підзвітність мають бути покладені в основу AI-системи правового забезпечення кадрового менеджменту.

Проведений огляд показує, що наукова література достатньо ґрунтовно розкриває три напрями: застосування ШІ в HR-менеджменті, використання нечіткої логіки для оцінювання персоналу та розроблення принципів відповідального застосування AI. Водночас більшість Fuzzy-HRM досліджень орієнтована на оцінювання кандидатів або працівників, а більшість AI-HRM публікацій – на ефективність, автоматизацію та прогнозування. Водночас недостатньо розробленими залишаються моделі, у яких ШІ та нечітка логіка застосовуються для оцінювання правового ризику кадрового рішення, особливо в органах військового управління. Це обґрунтовує доцільність розроблення AI-Fuzzy моделі, що поєднує NLP, RAG, граф знань, нечіткі правила Mamdani, пояснюваність і Human-in-the-Loop.

Методологія дослідження

Об'єктом дослідження є процес правового забезпечення кадрового менеджменту, а предметом дослідження – гібридна модель застосування ШІ та нечіткої логіки для оцінювання правового ризику кадрових рішень.

Кадрове рішення розглядається як множина взаємопов'язаних параметрів:

$$K = \{D, P, A, H, G, N, E\},$$

- де D – повнота документів;
 P – дотримання процедури;
 A – підтвердження повноважень;
 H – непорушення прав особи;
 G – захист персональних даних;
 N – релевантні нормативні підстави;
 E – експертна або AI-інтерпретація кадрової ситуації.

У загальному вигляді кадрове рішення можна вважати правомірним, якщо виконуються ключові правові умови:

$$\text{LegalDecision}(x) = \text{Authority}(x) \wedge \text{Grounds}(x) \wedge \text{Procedure}(x) \wedge \text{Documents}(x) \wedge \text{Rights}(x) \wedge \text{DataProtection}(x)$$

- де $\text{Authority}(x)$ – наявність повноважень у суб'єкта прийняття рішення;
 $\text{Grounds}(x)$ – наявність належних правових підстав;
 $\text{Procedure}(x)$ – дотримання встановленої процедури;

- Documents(x) – повнота документального оформлення;
 Rights(x) – непорушення прав особи;
 DataProtection(x) – дотримання вимог захисту персональних даних.

Якщо хоча б одна з умов виконується не повністю, виникає правовий ризик:

$$\text{LegalDecision}(x) \Rightarrow \text{LegalRisk}(x).$$

Однак у реальних кадрових процедурах кожна умова може мати частковий ступінь виконання:

$$\text{Authority}(x), \text{Grounds}(x), \text{Procedure}(x), \text{Documents}(x), \text{Rights}(x), \text{DataProtection}(x) \in [0;1].$$

Тому правовий ризик доцільно оцінювати не як бінарний стан, а як нечітку величину.

У роботі Zadeh L. A. [16], яка є фундаментальною для всієї теорії нечітких множин, запропоновано відмовитися від жорсткої бінарної логіки належності об'єкта до множини та замінити її поняттям ступеня належності. У класичній логіці елемент або належить до множини, або не належить. У нечіткій логіці він може належати до неї частково.

Правовий стан кадрового рішення також не завжди можна оцінити бінарно: «правомірне/неправомірне». Документи можуть бути частково повними, процедура - частково дотриманою, правова підстава – наявною, але недостатньо конкретизованою. Саме тому підхід Zadeh L. A. дозволяє подати правовий ризик як нечітку величину: $\mu_{\text{LegalRisk}}(x) \in [0;1]$.

Такий підхід створює теоретичну для оцінювання проміжних станів правової відповідності. У контексті кадрового менеджменту в органах військового управління це особливо важливо, оскільки кадрові рішення часто приймаються в умовах неповноти інформації, дефіциту часу та складного нормативного регулювання.

У запропонованій моделі штучний інтелект виконує роль аналітичного посередника між кадровими документами, нормативною базою і нечіткою системою оцінювання ризику і виконує п'ять основних функцій.

Перша функція ШІ – вилучення ознак із документів. Її виконує NLP-модуль, який аналізує документ і вилучає суттєві ознаки:

$$F_{doc} = AI_{NLP}(\text{Text}_{doc}),$$

де F_{doc} – набір вилучених ознак: дата, посада, підстава, суб'єкт рішення, додатки, посилання на норму, персональні дані.

NLP-модуль аналізує текст проєкту кадрового рішення, подання, рапорту, службової характеристики, довідки, висновку або іншого документа і виокремлює суттєві елементи: дату, суб'єкта рішення, вид кадрової процедури, посаду, підставу, наявність додатків, згадки про персональні дані, погодження, обмеження та потенційні суперечності.

Друга функція – пошук нормативних підстав. Її виконує RAG-модуль, який здійснює пошук релевантних правових норм:

$$N_q = AI_{RAG}(Q, DB_{law}),$$

де Q – запит або опис кадрової ситуації;

DB_{law} – база нормативних актів;

N_q – релевантні правові положення.

RAG-підхід важливий тому, що він дозволяє зменшити ризик необґрунтованої генерації відповіді та прив'язати рекомендацію до нормативного контексту.

Третя функція – класифікація типу кадрової процедури. Модель машинного навчання або LLM-класифікатор визначає тип рішення:

$$C = AI_{class}(F_{doc}),$$

де C – тип рішення.

Класифікатор може визначати, чи йдеться про призначення, переведення, звільнення, ротацию, направлення на навчання, дисциплінарне рішення або зміну умов проходження служби чи роботи.

Четверта функція – попереднє оцінювання параметрів нечіткої моделі. ШІ допомагає перетворити текстові й документальні дані у числові оцінки:

$$(D, P, A, H, G) = AI_{score}(F_{doc}, N_g, Rules).$$

Наприклад, якщо відсутній один із ключових додатків, система може знизити D ; якщо не виявлено погодження, знижується P ; якщо правова підстава не конкретизована, знижується показник процедури або документів.

П'ята функція – пояснення результату. LLM-модуль формує пояснювану рекомендацію:

$$Explanation = AI_{LLM}(R_{legal}, Factors, N_q).$$

Таким чином, ШІ не приймає остаточного рішення, а формує аналітичний вхід для нечіткої системи та пояснення для людини. Пояснення має бути зрозумілим для кадрового органу, юридичної служби або керівника. Наприклад: «ризик зумовлений неповним пакетом документів, нечітким посиланням на правову підставу та недостатнім контролем доступу до персональних даних».

Запропонована система має таку структуру:

$$Input \rightarrow AI_{NLP} \rightarrow AI_{RAG} \rightarrow KnowledgeGraph \rightarrow FuzzyInference \rightarrow XAI \rightarrow Human Review.$$

Таблиця 1 подає функціональні модулі системи.

Таблиця 1: Архітектура AI-Fuzzy системи оцінювання правових ризиків кадрових рішень

Модуль	Призначення	Результат
Модуль введення даних (<i>Input</i>)	Завантаження проекту кадрового рішення та документів	Текстові й структуровані дані
NLP-модуль (AI_{NLP})	Вилучення суттєвих кадрово-правових ознак	Ознаки F_{doc}
RAG-модуль (AI_{RAG})	Пошук релевантних норм у нормативній базі	Набір норм N_q
Граф знань (<i>KnowledgeGraph</i>)	Зв'язування особи, посади, процедури, документа, норми та ризику	Контекст кадрового рішення
Експертна база правил, Нечітка система Mamdani (<i>FuzzyInference</i>)	Формалізація правових умов, оцінювання рівня правового ризику	Правила перевірки R_{legal}
ML-модуль, XAI-модуль (<i>XAI</i>)	Калібрування ваг, порогів і функцій належності, пояснення чинників ризику	Адаптовані параметри, пояснюваний висновок
Human-in-the-Loop (<i>Human Review</i>)	Перевірка людиною	Остаточне кадрове рішення

Джерело: розроблено автором

Цей процес є послідовним, але не механічним. На кожному етапі можливе уточнення даних.

Якщо система не знаходить правової підстави, вона не повинна автоматично робити висновок про незаконність рішення. Вона має зафіксувати ризик, пояснити його і передати результат людині. Такий підхід відповідає принципу:

$$AI \neq FinalDecision,$$

$$AI = DecisionSupport.$$

Тобто ШІ виконує аналітичну, пошукову, пояснювальну і прогностичну функції, але не замінює уповноважену посадову особу або юридичну службу.

Результати

Вхідними змінними нечіткої моделі є (таблиця 2):

$$X = \{D, P, A, H, G\}$$

Таблиця 2: Вхідні змінні нечіткої моделі

Позначення	Назва	Зміст	Джерело оцінювання
<i>D</i>	<i>Documents</i>	повнота документів	AI-аналіз комплексу документів
<i>P</i>	<i>Procedure</i>	дотримання процедури	експертні правила + RAG
<i>A</i>	<i>Authority</i>	підтвердженість повноважень	граф знань + нормативна база
<i>H</i>	<i>HumanRights</i>	непорушення прав особи	правовий аналіз + LLM-пояснення
<i>G</i>	<i>DataProtection</i>	захист персональних даних	аудит доступу + правила безпеки

Джерело: розроблено автором.

Усі змінні нормуються в інтервалі:

$$x_i \in [0; 100],$$

де 0 – означає повну невідповідність;
100 – повну відповідність.

Вихідною змінною є інтегральний правовий ризик: $Y = R_{Legal}$, де $R_{Legal} \in [0; 100]$.

А для нормованого значення інтегрального правового ризику маємо вираз:

$$R_{Legal}^{norm} = \frac{R_{Legal}}{100}.$$

Для інтерпретації результатів оцінювання нормованого значення інтегрального правового ризику введемо шкалу, яку наведено у таблиці 3:

Таблиця 3. Шкала інтерпретації

Інтервал	Рівень ризику	Управлінське значення
0–0.33	Низький	можливе подання на затвердження
0.33–0.66	Середній	потрібна додаткова юридична перевірка
0.66–1.00	Високий	рішення потребує доопрацювання або зупинення

Джерело: розроблено автором.

Шкала інтерпретації нормованого значення інтегрального правового ризику побудована на основі поділу інтервалу $[0;1]$ на три рівні зони, що відповідають лінгвістичним термам нечіткої моделі: *Low*, *Medium*, *High*. Межі 0.33 і 0.66 визначені як наближені значення однієї та двох третин нормованого інтервалу ризику.

Такий поділ забезпечує просту управлінську інтерпретацію результату: низький ризик означає можливість подання рішення на затвердження, середній ризик — потребу в додатковій юридичній перевірці, високий ризик — необхідність доопрацювання або зупинення кадрового рішення.

Шкала має експертно-інтерпретаційний характер і може уточнюватися після накопичення емпіричних даних. У подальших дослідженнях межі шкали можна уточнити через експертне опитування військових юристів, або аналіз скасованих кадрових наказів чи скарг і судових спорів, або порівняння результатів моделі з фактичними наслідками рішень тощо.

Оскільки більше значення відповідає більшій правовій відповідності, то: *Low* означає низьку відповідність та високий ризик, *Medium* – часткову відповідність та середній ризик, а *High* – високу відповідність та низький ризик.

При цьому, для рівнів правової відповідності введемо наступні функції:
низький рівень відповідності:

$$\mu_{Low}(x) = \begin{cases} 1, & 0 \leq x \leq 30, \\ \frac{50-x}{20}, & 30 < x < 50 \\ 0, & x \geq 50 \end{cases}$$

середній рівень відповідності:

$$\mu_{Medium}(x) = \begin{cases} 0, & x \leq 30 \\ \frac{50-x}{20}, & 30 < x < 50 \\ 1, & 50 \leq x \leq 70 \\ \frac{90-x}{20}, & 70 < x < 90 \\ 0, & x \geq 90 \end{cases}$$

високий рівень відповідності:

$$\mu_{High}(x) = \begin{cases} 0, & x \leq 70, \\ \frac{x-70}{20}, & 70 < x < 90 \\ 1, & 90 \leq x \leq 100 \end{cases}$$

Для вихідної змінної R_{Legal} використовуються терми:

$$Risk = \{Low, Medium, High\},$$

з центрами:

$$C_{Low} = 20, C_{Medium} = 50, C_{High} = 80.$$

Mamdani E. H. та Assilian S. [17] показали, як евристичні правила, сформульовані людиною природною мовою, можна перетворити на формалізований механізм нечіткого висновку. Тобто логіку кадрового юриста або експерта з правового забезпечення можна перевести у правила типу:

$$IF Document = Low \wedge Procedure = Low THEN LegalRisk = High.$$

Такий підхід дозволяє не лише обчислювати ризик, а й пояснювати його причини. Крім того, експертні юридичні правила перетворюються на алгоритм попередньої правової діагностики кадрового рішення.

База правил Mamdani є експертно-нормативним відображення логіки правової перевірки кадрового рішення. Тобто правила не є довільними, а впливають із самої структури правомірності кадрового рішення, де оцінюються документи, процедура, повноваження, права особи та персональні дані. Це показники моделі (1).

База правил Mamdani має відображати типові правові ситуації, які виникають під час підготовки кадрового рішення. У правовому забезпеченні кадрового менеджменту ключовим є не лише наявність формального документа, а сукупна відповідність рішення умові (2).

Якщо хоча б одна з цих умов істотно не виконується, виникає правовий ризик (3).

Саме тому правила Mamdani побудовані навколо п'яти найбільш значущих груп факторів: документальної повноти, процедурної відповідності, повноважень суб'єкта, дотримання прав особи та захисту персональних даних.

База правил Mamdani має такий вигляд:

Правило 1: IF D = Low OR P = Low THEN Risk = High.

Сутність правила: якщо документи неповні або процедура суттєво порушена, кадрове рішення має високий правовий ризик, оскільки відсутня належна документальна чи процедурна основа

Правило 2: IF D = Medium AND P = Medium THEN Risk = Medium.

Сутність правила: якщо документи та процедура мають лише часткову відповідність, ризик оцінюється як середній, оскільки рішення не є явно неправомірним, але потребує додаткової юридичної перевірки.

Правило 3: IF A = High AND H = High AND G = High THEN Risk = Low.

Сутність правила: якщо повноваження підтверджені, права особи дотримані, а персональні дані захищені, правовий ризик є низьким.

Правило 4: IF D = High AND P = High AND A = High THEN Risk = Low.

Сутність правила: якщо документи повні, процедура дотримана, а посадова особа має належні повноваження, кадрове рішення має низький правовий ризик.

Правило 5: IF H = Medium OR G = Medium THEN Risk = Medium.

Сутність правила: якщо є середній ризик щодо прав особи або персональних даних, загальний ризик не можна вважати низьким, навіть якщо інші показники є прийнятними.

Правило 6: IF D = Medium AND A = High AND P = Medium THEN Risk = Medium.

Сутність правила: навіть за підтверджених повноважень середній рівень документів і процедури залишає кадрове рішення у зоні середнього правового ризику.

Для розширеної AI-Fuzzy моделі додати правила, пов'язані з результатами ШІ-аналізу.

Правило 7: IF AIConfidence = Low THEN Risk = Medium.

Сутність правила: якщо ШІ має низьку впевненість у результаті аналізу *AIConfidence*, рішення потребує додаткової перевірки людиною.

Правило 8: IF NormativeGrounds = Unclear THEN Risk = Medium.

Сутність правила: якщо нормативні підстави кадрового рішення *NormativeGrounds* нечіткі або не повністю ідентифіковані *Unclear*, ризик оцінюється як середній.

Правило 9: IF DocumentContradiction = Detected THEN Risk = High.

Сутність правила: якщо ШІ виявив *Detected* суперечності між документами *DocumentContradiction*, проектом наказу або додатками, правовий ризик є високим.

Правило 10: IF Appeal $\rightarrow$ P $\rightarrow$ rediction = High THEN Risk = High.

Сутність правила: якщо система прогнозує високу ймовірність скарги або оскарження кадрового рішення *AppealPrediction*, ризик оцінюється як високий.

Ці правила дозволяють враховувати не лише формальні кадрові показники, а й результати інтелектуального аналізу документів, нормативних джерел і прогнозування ризику оскарження.

На практиці значення *D, P, A, H* та *G* може визначатися не вручну, а за допомогою ШІ. Наприклад:

$$D = 100 - \frac{M_{miss}}{M_{reg}} \cdot 100,$$

де M_{miss} – кількість відсутніх обов'язкових документів;

M_{reg} – загальна кількість необхідних документів.

Показник процедури може оцінюватися так:

$$P = 100 - \sum_{k=1}^m v_k p_k,$$

де p_k – виявлене k -те процедурне порушення;
 v_k – вагомість k -го процедурного порушення.
 Показник повноважень A :

$$A = AI_{auth}(Subject, DecisionType, Norms),$$

де $Subject$ – суб'єкт, який приймає або підписує кадрове рішення;
 $DecisionType$ – тип кадрового рішення;
 $Norms$ – релевантні нормативно-правові акти, положення, накази, статuti або посадові інструкції;
 AI_{auth} – AI-модуль перевірки повноважень.

Ця формула означає, що AI-модуль перевірки повноважень AI_{auth} перевіряє, чи має відповідний суб'єкт право приймати саме цей тип кадрового рішення.

Показник непорушення прав H :

$$H = AI_{rights}(Decision, PersonnelData, LegalGuarantees).$$

де $Decision$ – зміст кадрового рішення
 $PersonnelData$ – відомості про військовослужбовця або працівника;
 $LegalGuarantees$ – правові гарантії, які повинні бути враховані;
 AI_{rights} – AI-модуль оцінювання ризику порушення прав особи.

Ця формула означає, що AI-модуль оцінювання ризику порушення прав особи AI_{rights} перевіряє, чи не порушує кадрове рішення права, гарантії та законні інтереси особи.

Показник захисту персональних даних G :

$$G = AI_{data}(AccessLogs, DataCategories, SecurityRules).$$

де $AccessLogs$ – журнали доступу до кадрової інформації;
 $DataCategories$ – категорії персональних даних, що обробляються;
 $SecurityRules$ – правила інформаційної безпеки, доступу, зберігання та обробки даних;
 AI_{data} – AI-модуль перевірки дотримання вимог щодо персональних даних.

Ця формула означає, що AI-модуль перевірки дотримання вимог щодо персональних даних AI_{data} перевіряє, чи належно захищені персональні дані, які використовуються під час підготовки, погодження або прийняття кадрового рішення.

Таким чином, ШІ перетворює складні документи та процедури у числові оцінки, які далі подаються до нечіткої системи.

Приклад розрахунку

1. Базовий сценарій

Нехай AI-модуль після аналізу проекту кадрового рішення визначив такі значення:

$$D = 55, P = 65, A = 90, H = 80, G = 70.$$

Для заданих значень отримані такі результати:

ступені належності:

значення $D = Medium(1.0)$ означає, що повнота документів перебуває у середній зоні: документи не є критично неповними, але їхній рівень не дозволяє вважати рішення повністю юридично безпечним;

значення $P = Medium(1.0)$ означає, що процедура також має середній рівень відповідності. Вона не є повністю порушеною, але потребує додаткової перевірки;

значення $A = High(0.5)$ означає, що повноваження суб'єкта кадрового рішення підтверджені на високому рівні;

значення $H = Medium(0.5) \wedge High(0.5)$ означає, що показник непорушення прав особи перебуває на межі між середнім і високим рівнями, тобто права загалом враховані, але певна правова невизначеність зберігається;

значення $G = Medium(1.0)$ означає, що захист персональних даних оцінюється як середній. Саме цей показник не дозволяє класифікувати рішення як низько ризикове;

активація правил:

$\alpha_1 = \max(0,0) = 0$ – перше правило не активується, оскільки ні документи, ні процедура не мають низького рівня;

$\alpha_2 = \min(1,1) = 1$ – друге правило активується повністю, оскільки і документи, і процедура мають середній рівень, що формує середній правовий ризик;

$\alpha_3 = \min(1,0.5,0) = 0$ – третє правило не активується, оскільки для низького ризику необхідно, щоб повноваження, права особи та персональні дані одночасно мали високий рівень. Тобто G не має високої належності;

$\alpha_4 = \min(0,0,1) = 0$ – четверте правило також не активується, бо документи й процедура не належать до високого рівня;

$\alpha_5 = \max(0.5,1) = 1$ – п'яте правило активується повністю, оскільки персональні дані мають середній рівень, а права особи частково належать до середнього рівня, що відповідає середньому ризику;

$\alpha_6 = \min(1,1,1) = 1$ – шосте правило активується повністю, оскільки документи та процедура мають середній рівень, а повноваження – високий. Це підтверджує, що навіть за належних повноважень середній рівень документів і процедури залишає рішення у зоні середнього ризику.

Отже, агрегований результат має вигляд:

$$\alpha_{Low} = 0, \alpha_{Medium} = 1, \alpha_{High} = 0.$$

Тобто в системі нечіткого висновку домінує саме середній рівень правового ризику.

Після деафазифікації отримаємо:

$$R^* = \frac{\alpha_{Low}C_{Low} + \alpha_{Medium}C_{Medium} + \alpha_{High}C_{High}}{\alpha_{Low} + \alpha_{Medium} + \alpha_{High}}$$

$$R^* = \frac{0 \cdot 20 + 1 \cdot 50 + 0 \cdot 80}{0 \cdot 20 + 1 \cdot 50 + 0 \cdot 80} = \frac{50}{50} = 1$$

$$R_{legal}^{norm} = 0.50$$

Отже:

$$R_{legal} = 0.50.$$

Це відповідає **середньому правовому ризику**.

Отже, кадрові рішення не має ознак критичного правового ризику, але й не може бути віднесене до безпечних. Його доцільно направити на додаткову юридичну перевірку, зокрема щодо повноти документів, дотримання процедури та режиму захисту персональних даних.

2. Покращений сценарій

Після рекомендацій ШІ кадровий орган доукомплектував документи, уточнив правову підставу та перевірів процедуру. Нові значення:

$$D = 75, P = 72, A = 90, H = 82, G = 78.$$

Порівняно з базовим сценарієм покращилися насамперед: $D: 55 \rightarrow 75$, $P: 65 \rightarrow 72$, $G: 70 \rightarrow 78$. Тобто кадрові рішення стало більш документально забезпеченим, процедурно уточненим і краще захищеним з погляду персональних даних.

Ступені належності мають такі значення:

$\mu_{DMedium}(75) = \frac{90-75}{20} = 0.75$ та $\mu_{DHigh}(75) = \frac{75-70}{20} = 0.25$ – документи вже не перебувають лише в середній зоні, а частково переходять до високого рівня відповідності, але ще не досягають повністю високого рівня;

$\mu_{PMedium}(72) = \frac{90-72}{20} = 0.90$ та $\mu_{PHigh}(72) = \frac{72-70}{20} = 0.10$ – процедура переважно ще належить до середнього рівня, але вже з'являється початкова належність до високого рівня;

$\mu_{AHigh}(90) = 1.0$ – повноваження суб'єкта кадрового рішення повністю підтверджені;

$\mu_{HMedium}(82) = \frac{90-82}{20} = 0.40$ та $\mu_{HHigh}(82) = \frac{82-70}{20} = 0.60$ – показник прав особи вже більше тяжіє до високого рівня, хоча залишкова середня належність ще зберігається;

$\mu_{GMedium}(78) = \frac{90-78}{20} = 0.60$ та $\mu_{GHigh}(78) = \frac{78-70}{20} = 0.40$ – захист персональних даних покращився, але ще не є повністю високим.

Активовані правила:

$\alpha_2 = \min(0.75, 0.90) = 0.75$ – правило підтримує середній ризик, оскільки документи й процедура ще мають суттєву належність до середнього рівня;

$\alpha_3 = \min(1.0, 0.60, 0.40) = 0.40$ – правило підтримує низький ризик, оскільки повноваження, права особи та персональні дані частково перебувають у високій зоні. Але через те, що $G = 78$ має лише 0.40 належності до високого рівня, сила цього правила обмежена;

$\alpha_4 = \min(0.25, 0.10, 1.0) = 0.10$ – правило також підтримує низький ризик, але слабо, тому що документи й процедура ще мають невисоку належність до високого рівня;

$\alpha_5 = \max(0.40, 0.60) = 0.60$ – правило підтримує середній ризик, оскільки права особи та персональні дані ще частково залишаються в середній зоні.

Після агрегування активованих правил отримуємо такі значення:

$$\alpha_{Low} = \max(0.40, 0.10) = 0.40,$$

$$\alpha_{Medium} = \max(0.75, 0.60) = 0.75,$$

$$\alpha_{High} = 0.$$

Тобто високий ризик уже не активується, а середній ризик залишається домінантним, хоча низький ризик також починає проявлятися.

Дефазифікація з центрами $C_{Low} = 20$, $C_{Medium} = 50$ та $C_{High} = 80$ дає такі результати.

$$R^* = \frac{0.40 \cdot 20 + 0.75 \cdot 50 + 0.80}{0.40 + 0.75}.$$

$$R^* = \frac{8 + 37.5}{1.15} = 39.57.$$

$$R_{legal}^{norm} = 0.396.$$

Отже:

$$R_{legal} = 0.396.$$

Це відповідає нижній межі середнього ризику.

Таким чином, кадрові рішення стало значно безпечнішим порівняно з базовим сценарієм, де ризик становив $R_{legal}^{norm} = 0.50$. Проте рішення ще не досягло низького ризику, оскільки документи, процедура та захист персональних даних залишаються частково в середній зоні. Тобто рішення може бути подане на подальше погодження, але бажано додатково уточнити окремі елементи документального оформлення, процедури та режиму обробки персональних даних.

3. Порівняння сценаріїв

Таблиця 3: Порівняння базового та покращеного сценаріїв

Показник	Базовий сценарій	Покращений сценарій
Повнота документів D	55	75
Дотримання процедури P	65	72

Показник	Базовий сценарій	Покращений сценарій
Повноваження А	90	90
Захист прав Н	80	82
Персональні дані G	70	78
Нечіткий ризик R_{legal}	0.50	0.396
Рівень ризику	середній	середній, ближче до низького
Рекомендація	додаткова перевірка обов'язкова	можливе подання після уточнень

Джерело: розроблено автором.

Обговорення

Результати демонструють, що поєднання ШІ та нечіткої логіки дозволяє отримати більш обережну і пояснювану оцінку правового ризику кадрового рішення. У базовому сценарії $R_{legal}=0.50$, тобто рішення не належить до високоризикових, але потребує додаткової юридичної перевірки. Основними причинами є середній рівень повноти документів, середній рівень дотримання процедури та недостатньо високий рівень захисту персональних даних.

Після покращення документального забезпечення та процедури ризик зменшується до $R_{legal}=0.396$. Проте модель не класифікує його як повністю низький, оскільки окремі показники все ще залишаються в зоні середньої відповідності. Це свідчить про те, що нечітка система виконує роль «обережного правового фільтра».

Важливою перевагою запропонованого підходу є те, що ШІ не лише обчислює ризик, а й пояснює його походження. Наприклад, система може сформулювати висновок: “Ризик зумовлений середньою повнотою документів, відсутністю чіткої фіксації процедури погодження та недостатнім рівнем контролю доступу до персональних даних”. Така пояснюваність є критично важливою для довіри до системи.

Порівняно зі звичайною зваженою моделлю:

$$R_{legal} = \sum_{i=1}^n w_i R_i,$$

AI-Fuzzy модель має суттєву перевагу: вона не дозволяє автоматично компенсувати процедурні порушення іншими позитивними показниками. Наприклад, навіть за високого рівня повноважень середній рівень документального забезпечення може залишати ризик у середній зоні.

Водночас застосування ШІ у кадровому менеджменті має бути обмежене принципом людського контролю. Це зазначає, що ШІ-інструменти не можуть бути авторами наукових публікацій, оскільки не несуть відповідальності за зміст; аналогічно в кадровому менеджменті ШІ не може нести відповідальність за правові наслідки кадрового рішення – вона має залишатися за уповноваженою людиною.

Наукова новизна

Наукова новизна дослідження полягає в такому:

запропоновано гібридну AI-Fuzzy модель оцінювання правового ризику кадрових рішень;

обґрунтовано роль ШІ як інструменту вилучення ознак, пошуку нормативних підстав, аналізу документів, формування пояснень і калібрування нечіткої моделі;

формалізовано правовий ризик кадрового рішення як нечітку величину, що дозволяє враховувати проміжні стани правової відповідності;

розроблено базу правил Mamdani для оцінювання правового ризику кадрового рішення;

показано, що AI-Fuzzy модель є більш обережною порівняно з простою зваженою моделлю, оскільки не допускає повної компенсації документальних або процедурних недоліків іншими позитивними показниками;

запропоновано архітектуру інтелектуальної системи підтримки правового забезпечення кадрового менеджменту з модулями NLP, RAG, графа знань, нечіткого висновку, XAI та Human-in-the-Loop.

Практична значущість

Практична значущість запропонованої моделі полягає в можливості її використання для: попередньої перевірки кадрових наказів;

оцінювання правового ризику призначення на посаду;

аналізу процедур переведення, звільнення або ротації;

контролю повноти кадрових документів;

перевірки повноважень посадових осіб;

виявлення ризиків порушення прав особи;

перевірки режиму обробки персональних даних;

формування пояснюваних рекомендацій для кадрового органу та юридичної служби;

побудови аналітичної панелі кадрово-правових ризиків;

моніторингу динаміки ризиків у кадровій системі.

У практичному впровадженні система може працювати так:

Documents → *AI Extraction* → *Legal R* ↔ *etriel* → *Fuzzy Risk* → *Explanation*
→ *Human Decision*

Тобто кадровий документ (*Documents*) автоматично аналізується ШІ (*AI Extraction*), релевантні норми (*Legal R* ↔ *etriel*) знаходяться через RAG-модуль, ризик обчислюється нечіткою системою (*Fuzzy Risk*), а результат (*Explanation*) передається людині для остаточного рішення (*Human Decision*).

Обмеження дослідження

Запропонована модель має концептуально-методичний характер і потребує подальшої емпіричної перевірки. Основні обмеження полягають у такому:

функції належності задано експертно і вони потребують валідації на реальних кадрових кейсах;

база правил Mamdani є демонстраційною і має бути розширена для різних типів кадрових рішень;

AI-модулі можуть помилятися під час вилучення ознак або тлумачення документів;

якість результату залежить від актуальності нормативної бази;

модель не враховує всі спеціальні правові режими, винятки та процедурні особливості;

існує ризик алгоритмічної упередженості при навчанні ML-модуля на історичних кадрових рішеннях;

система потребує аудиту, журналювання, контролю доступу та механізмів пояснюваності;

модель не може замінювати висновок юридичної служби або рішення уповноваженої посадової особи.

Висновки

У статті розроблено гібридну модель застосування штучного інтелекту та нечіткої логіки для оцінювання правових ризиків кадрових рішень у системі правового забезпечення кадрового менеджменту. Обґрунтовано, що правовий ризик кадрового рішення доцільно розглядати як нечітку категорію, оскільки реальні кадрові процедури часто характеризуються частковою

відповідністю документів, неповною процедурною визначеністю, потребою уточнення правових підстав і необхідністю контролю персональних даних.

Запропонована AI-Fuzzy модель включає NLP-модуль, RAG-модуль, граф знань, експертну базу правил, нечітку систему Mamdani, модуль машинного навчання, ХAI-модуль та блок Human-in-the-Loop. ШІ використовується для вилучення ознак, пошуку нормативних підстав, аналізу документів, попереднього оцінювання змінних і пояснення результатів, а нечітка логіка – для розрахунку інтегрального правового ризику.

На прикладі умовного кадрового рішення показано, що базовий сценарій дає: $R_{legal} = 0.50$, що відповідає середньому правовому ризику. Після покращення документального забезпечення та процедури ризик знижується до: $R_{legal} = 0.396$.

Це свідчить, що модель дозволяє не лише оцінити рівень ризику, а й визначити чинники, які потребують доопрацювання.

Перспективами подальших досліджень є створення програмного прототипу AI-Fuzzy системи, розширення бази нечітких правил, формування датасету кадрово-правових кейсів, інтеграція з нормативною базою, розроблення графа знань кадрових процедур, тестування точності AI-модулів, перевірка ризиків алгоритмічної упередженості та створення аналітичної панелі для кадрових і юридичних підрозділів.

Фінансування

Це дослідження не отримало конкретної фінансової підтримки.

Конкуруючі інтереси

Автори заявляють, що у них немає конкуруючих інтересів.

Використання штучного інтелекту

Інструменти штучного інтелекту використовувалися для пошуку публікацій з питань правового забезпечення кадрового менеджменту у Збройних Силах України та визначення факторів, які впливають на формування потреби у військових юристах, а також методів та методик визначення потреби у військових юристах.

Список використаних джерел

1. Scopus content policy and selection [Електронний ресурс] / Elsevier. URL: <https://www.elsevier.com/products/scopus/content/content-policy-and-selection> (дата звернення: 08.07.2026).
2. AI Principles [Електронний ресурс] / Organisation for Economic Co-operation and Development. 2019, updated 2024. URL: <https://www.oecd.org/en/topics/sub-issues/ai-principles.html> (дата звернення: 08.07.2026).
3. Artificial Intelligence Risk Management Framework [Електронний ресурс] / National Institute of Standards and Technology. 2023. URL: <https://www.nist.gov/> (дата звернення: 08.07.2026).
4. European Union. Artificial Intelligence Act. Regulation (EU) 2024/1689. URL: <https://eur-lex.europa.eu/legal-content/EN/TXT> (дата звернення: 08.07.2026).
5. Artificial Intelligence Act. Annex III: High-Risk AI Systems [Електронний ресурс] / European Union. URL: <https://artificialintelligenceact.eu/annex/3/> (дата звернення: 08.07.2026).
6. Shaout A., Al-Shammari M. Fuzzy logic modeling for performance appraisal systems: A framework for empirical evaluation, Expert Systems with Applications, Volume 14, Issue 3, 1998, Pages 323-328, ISSN 0957-4174, URL: [https://doi.org/10.1016/S0957-4174\(97\)00085-7](https://doi.org/10.1016/S0957-4174(97)00085-7). (<https://www.sciencedirect.com/science/article/pii/S0957417497000857>).

7. Sun C. C. A performance evaluation model by integrating fuzzy AHP and fuzzy TOPSIS methods, *Expert Systems with Applications*, Volume 37, Issue 12, 2010, Pages 7745-7754, ISSN 0957-4174, URL: <https://doi.org/10.1016/j.eswa.2010.04.066>. (<https://www.sciencedirect.com/science/article/pii/S0957417410003660>).
8. Dursun M., Karsak E. E. A fuzzy MCDM approach for personnel selection. *Expert Systems with Applications*. 2010. 37. 4324-4330. DOI: <https://doi.org/10.1016/j.eswa.2009.11.067>. URL: https://www.researchgate.net/publication/220216329_A_fuzzy_MCDM_approach_for_personnel_selection.
9. Kusumawardani R. P., Agintiara M. Application of Fuzzy AHP-TOPSIS Method for Decision Making in Human Resource Manager Selection Process. *Procedia Computer Science*. 72. 638-646. DOI: <https://doi.org/10.1016/j.procs.2015.12.173>. URL: https://www.researchgate.net/publication/289991287_Application_of_Fuzzy_AHP-TOPSIS_Method_for_Decision_Making_in_Human_Resource_Manager_Selection_Process/
10. Hoque S. et al. When LLM meets Fuzzy-TOPSIS for Personnel Selection through Automated Profile Analysis. *Institute of Electrical and Electronics Engineers (IEEE)*. 2026. 14, pp. 15778–15794 ISSN 2169-3536, DOI: <https://doi.org/10.1109/access.2026.3658575>, URL: <http://dx.doi.org/10.1109/ACCESS.2026.3658575>.
11. Barocas, Solon and Selbst, Andrew D., *Big Data's Disparate Impact* (2016). 104 *California Law Review* 671 (2016), Available at SSRN: <https://ssrn.com/abstract=2477899> or <http://dx.doi.org/10.2139/ssrn.2477899>.
12. OECD. *AI Principles*. 2019; updated 2024. URL: <https://www.oecd.org/en/topics/ai-principles.html>.
13. Binns, Reuben, *Human Judgement in Algorithmic Loops; Individual Justice and Automated Decision-Making* (September 11, 2019). Available at SSRN: <https://ssrn.com/abstract=3452030> or <http://dx.doi.org/10.2139/ssrn.3452030>.
14. NATO. Summary of NATO's Revised Artificial Intelligence Strategy. 2024. URL: <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2024/07/10/summary-of-natos-revised-artificial-intelligence-ai-strategy>.
15. NATO. Summary of the NATO Artificial Intelligence Strategy. 2021. URL: <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2021/10/22/summary-of-the-nato-artificial-intelligence-strategy>.
16. Zadeh L. A. Fuzzy sets. *Information and Control*. 1965. Vol. 8, Issue 3. P. 338–353. [https://doi.org/10.1016/S0019-9958\(65\)90241-X](https://doi.org/10.1016/S0019-9958(65)90241-X).
17. Mamdani E. H., Assilian S. An experiment in linguistic synthesis with a fuzzy logic controller. *International Journal of Man-Machine Studies*. 1975. Vol. 7, Issue 1. P. 1–13. [https://doi.org/10.1016/S0020-7373\(75\)80002-2](https://doi.org/10.1016/S0020-7373(75)80002-2).

References

1. Elsevier. (n.d.). *Scopus content policy and selection*. Retrieved July 8, 2026, from <https://www.elsevier.com/products/scopus/content/content-policy-and-selection>.
2. Organisation for Economic Co-operation and Development. (2024). *AI principles*. <https://www.oecd.org/en/topics/ai-principles.html>.
3. Tabassi, E. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)* (NIST AI 100-1). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.100-1>.
4. European Parliament & Council of the European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No. 300/2008, (EU) No. 167/2013, (EU) No. 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU)

- 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act). *Official Journal of the European Union*. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>.
5. European Union. (2024). *Artificial Intelligence Act: Annex III—High-risk AI systems*. <https://artificialintelligenceact.eu/annex/3/>
6. Shaout, A., & Al-Shammari, M. (1998). Fuzzy logic modeling for performance appraisal systems: A framework for empirical evaluation. *Expert Systems with Applications*, 14(3), 323–328. [https://doi.org/10.1016/S0957-4174\(97\)00085-7](https://doi.org/10.1016/S0957-4174(97)00085-7).
7. Sun, C.-C. (2010). A performance evaluation model by integrating fuzzy AHP and fuzzy TOPSIS methods. *Expert Systems with Applications*, 37(12), 7745–7754. <https://doi.org/10.1016/j.eswa.2010.04.066>.
8. Dursun, M., & Karsak, E. E. (2010). A fuzzy MCDM approach for personnel selection. *Expert Systems with Applications*, 37(6), 4324–4330. <https://doi.org/10.1016/j.eswa.2009.11.067>.
9. Kusumawardani, R. P., & Agintiara, M. (2015). Application of fuzzy AHP-TOPSIS method for decision making in human resource manager selection process. *Procedia Computer Science*, 72, 638–646. <https://doi.org/10.1016/j.procs.2015.12.173>.
10. Hoque, S., Karim, A. A. J., Alam, M. G. R., & Gope, N. (2026). When LLM meets Fuzzy-TOPSIS for personnel selection through automated profile analysis. *IEEE Access*, 14, 15778–15794. <https://doi.org/10.1109/ACCESS.2026.3658575>.
11. Barocas, S., & Selbst, A. D. (2016). Big data's disparate impact. *California Law Review*, 104(3), 671–732. <https://ssrn.com/abstract=2477899>.
12. Organisation for Economic Co-operation and Development. (2024). *AI principles*. <https://www.oecd.org/en/topics/ai-principles.html>.
13. Binns, R. (2019). *Human judgement in algorithmic loops: Individual justice and automated decision-making*. SSRN. <https://ssrn.com/abstract=3452030>.
14. North Atlantic Treaty Organization. (2024, July 10). *Summary of NATO's revised artificial intelligence (AI) strategy*. <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2024/07/10/summary-of-natos-revised-artificial-intelligence-ai-strategy>.
15. North Atlantic Treaty Organization. (2021, October 22). *Summary of the NATO artificial intelligence strategy*. <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2021/10/22/summary-of-the-nato-artificial-intelligence-strategy> [Belgium].
16. Zadeh, L. A. (1965). Fuzzy sets. *Information and Control*, 8(3), 338–353. [https://doi.org/10.1016/S0019-9958\(65\)90241-X](https://doi.org/10.1016/S0019-9958(65)90241-X).
17. Mamdani, E. H., & Assilian, S. (1975). An experiment in linguistic synthesis with a fuzzy logic controller. *International Journal of Man-Machine Studies*, 7(1), 1–13. [https://doi.org/10.1016/S0020-7373\(75\)80002-2](https://doi.org/10.1016/S0020-7373(75)80002-2).



This is an open access journal and all published articles are licensed under a Creative Commons «Attribution» 4.0.