

ІНФОРМАЦІЙНІ, УПРАВЛІНСЬКІ, ЕКОНОМІЧНІ АСПЕКТИ БЕЗПЕКИ ЖИТТЄДІЯЛЬНОСТІ ЛЮДИНИ У СУЧАСНИХ УМОВАХ

*Ялова К. А., студ. (гр. БС-22, ФБМІ КПІ ім. Ігоря Сікорського);
Мітюк Л. О., канд. техн. наук, доц. (каф. ОППЦБ КПІ ім. Ігоря Сікорського)*

Анотація. Розглянуто питання застосування технологій штучного інтелекту в мобільній медицині (mHealth) з точки зору інформаційних, управлінських та економічних аспектів безпеки життєдіяльності людини. На прикладі розробки AI-консультанта для діагностики шкірних захворювань проаналізовано ключові ризики: точність діагностики, надійність та захист інформації. Досліджено управлінські аспекти інтеграції таких застосунків у систему охорони здоров'я та їх економічний ефект для раннього виявлення загроз здоров'ю та зниження навантаження на медичні заклади. Запропоновано заходи для забезпечення надійності та безпеки AI-рішень у сфері охорони здоров'я.

Ключові слова: штучний інтелект, мобільна медицина (mHealth), безпека пацієнта, інформаційна безпека, медичні дані, управління ризиками, телемедицина.

Abstract. The issues of applying artificial intelligence technologies in mobile medicine (mHealth) from the perspective of informational, managerial, and economic aspects of human life safety are considered. Using the example of the development of an AI consultant for diagnosing skin diseases, key risks are analyzed: diagnostic accuracy, reliability, and information security. The managerial aspects of integrating such applications into the healthcare system and their economic effect on the early detection of health threats and reducing the burden on medical institutions are explored. Measures are proposed to ensure the reliability and safety of AI solutions in the healthcare sector.

Keywords: artificial intelligence, mobile medicine (mHealth), patient safety, information security, medical data, risk management, telemedicine.

Вступ. Стрімке впровадження технологій штучного інтелекту (AI) у повсякденне життя відкриває нові горизонти для сфери охорони здоров'я. Особливе місце займають мобільні застосунки, що використовують AI як консультантів для попередньої діагностики, наприклад, дерматологічних захворювань. Потенціал таких технологій величезний: вони обіцяють суттєво підвищити доступність медичної допомоги та сприяти ранньому виявленню небезпечних станів. Водночас, така глибока інтеграція AI в медицину нерозривно пов'язана з новими викликами для безпеки життєдіяльності людини. Виникає гостра необхідність у комплексному аналізі цих інновацій не лише з технічної, але й з безпекової точки зору.

В інформаційному аспекті, ключовими є надійність та точність діагнозів, згенерованих алгоритмом, а також гарантування належного захисту та конфіденційності надзвичайно чутливих медичних даних пацієнтів.

В управлінському аспекті, необхідно визначити місце цих інструментів в існуючій медичній практиці. Чи стануть вони сертифікованими помічниками лікаря, чи залишаться нерегульованими інструментами, що можуть призвести до самолікування та хибного відчуття безпеки? Окремо стоїть питання відповідальності за діагностичні помилки AI.

В економічному аспекті, слід збалансувати потенційну економію ресурсів від автоматизації та ранньої діагностики з можливими економічними збитками від наслідків невірних рішень AI, які можуть призвести до прогресування хвороби.

Аналіз стану питання. Проблема забезпечення безпеки життєдіяльності людини в умовах стрімкої цифровізації суспільства є предметом активних наукових дискусій. Останніми роками особлива увага приділяється інтеграції технологій штучного інтелекту у сфери, що безпосередньо впливають на здоров'я та добробут, особливо в мобільну медицину.

Мета роботи: дослідження зазначених інформаційних, управлінських та економічних аспектів безпеки, пов'язаних із використанням AI-консультантів у сучасній системі охорони здоров'я, для визначення шляхів мінімізації потенційних ризиків для здоров'я та добробуту людини.

Методики, матеріали і результати досліджень. Наразі стрімке впровадження штучного інтелекту в охорону здоров'я, зокрема у формі мобільних консультантів для діагностики захворювань, відкриває нову еру в медицині. Ця технологічна революція несе в собі величезний потенціал, але водночас створює складні, багатовимірні виклики для безпеки життєдіяльності людини. Ефективність цих інструментів має оцінюватись на основі багатьох факторів, адже необхідно розуміти, що їх безпечне впровадження вимагає комплексного аналізу на перетині інформаційних, управлінських та економічних аспектів.

Існують дослідження, які активно поширюють думки про те, що згортована нейронна мережа здатна класифікувати рак шкіри з рівнем компетенції, “порівняним з дерматологами” [1]. Така технічна спроможність виступає суперечливим явищем, що може спричиняти як прогрес, так і ризики. З одного боку, він підтвердив величезний потенціал технології, але з іншого – створив небезпечну громадську перцепцію та, як наслідок, ринковий попит на нерегульовані інструменти. В результаті наукове досягнення мимоволі перетворилося на інформаційний ризик – ілюзію компетентності, яка не враховує прірву між якісними даними в лабораторії та фотографіями низької якості, зробленими пацієнтом у реальному житті.

Трактування штучного інтелекту як лікаря вступає у конфлікт із традиційними управлінськими парадигмами. Вона стимулює появу “тіньової” системи охорони здоров'я, де пацієнти використовують несертифіковані

застосунки для самодіагностики, оминаючи традиційну медицину. Саме тому регулятори, як-от Європейська Комісія у своєму “Акті про штучний інтелект”, цілком слушно пропонують класифікувати медичні діагностичні системи як “високоризикові” [2]. Однак метою регулювання має бути не усунення технології, а формування механізмів, що забезпечують її безпечне та кероване застосування. Найбезпечніше управлінське рішення – це заборона на автономну діагностику та інтеграція AI виключно як системи підтримки прийняття рішень для кваліфікованого лікаря [3].

Очевидність потреби в посиленому управлінському контролі проявляється при розгляді економічних аспектів, які характеризуються істотною нерівномірністю розподілу ризиків. Потенційна індивідуальна вигода від використання застосунку, до прикладу економія часу та коштів на візиті до лікаря є привабливою, але відносно малою. На противагу цьому, потенційні суспільні збитки від однієї-єдиної хибнонегативної помилки (пропущеної меланоми) є катастрофічними. Вартість лікування задавненої хвороби лягає величезним тягарем як на самого пацієнта, так і на всю систему охорони здоров’я [4]. Ця економічна асиметрія доводить, що вільний ринок не здатний самостійно регулювати цю технологію, ціна помилки є занадто високою.

Фундаментальний виклик полягає не в окремому вдосконаленні інформаційної точності, управлінської політики чи економічної моделі. Він полягає в управлінні їхньою невід’ємною взаємозалежністю. Інформаційний збій, такий як алгоритмічна упередженість, неминуче провокує управлінську кризу (втрату довіри, судові позови) та економічну катастрофу (нерівність у наданні допомоги, марнування ресурсів). Отже, безпека життєдіяльності в епоху AI – це не характеристика самого алгоритму. Це якість та надійність всієї системи, що вибудована навколо нього, покликана стримувати цей каскадний ефект збою та гарантувати, що технологія служить людині, а не створює нові, невидимі загрози.

Висновки. Безпека AI-консультантів у медицині – це не питання одного лише коду чи точності алгоритму. Справжні виклики лежать на перетині технологій, управлінських рішень та економіки.

Лабораторні успіхи AI створюють небезпечну ілюзію. В реальному світі якість “домашніх” фотографій та алгоритмічна упередженість створюють інформаційні ризики, де ціна однієї помилки – це пропущений рак. Економічні наслідки хибного діагнозу є катастрофічними і для пацієнта, і для системи охорони здоров’я. Вони неспівмірні з будь-якою економією на візиті до лікаря.

Безпечна стратегія передбачає суворе регулювання та використання AI лише як допоміжного інструменту для прийняття рішень лікарем з відповідною кваліфікацією.

Література

1. Esteva A, Kuprel B, Novoa RA, Ko J, Swetter SM, Blau HM, Thrun S. Dermatologist-level classification of skin cancer with deep neural networks. *Nature*. 2017 Feb 2;542(7639):115-118. doi: 10.1038/nature21056. Epub 2017 Jan 25. Erratum in: *Nature*. 2017 Jun 28;546(7660):686. doi: 10.1038/nature22985. PMID: 28117445; PMCID: PMC8382232.
2. European Commission. (2022). "Proposal for a Regulation on Artificial Intelligence (Artificial Intelligence Act)".
3. Topol, E. J. (2019). "Deep Medicine: How Artificial Intelligence Can Make Healthcare Human Again".
4. MedTech Europe. (2020). "The Socio-Economic Impact of AI in Healthcare".