

1. Курко Я.В. Психофізіологічні особливості осіб, які займаються плаванням за різних типів погоди: автореф. дис... канд. мед. наук: 14.03.03. Львів, 2007. 22 с.

2. Комп'ютерна програма "Вимірювання простої слухомоторної реакції (Reaction-test)" / Курко Я.В. А.с. № 13683 від 20.07.2005. Державний департамент інтелектуальної власності Я.В; Заявл. 31.05.05; опубл. 30.04.06; Бюл № 8, серія KB № 6018.– С. 110-111.

УДК 004.9

Т. А. Липак, аспірант

Тернопільський національний технічний університет імені І. Пулюя, Україна

ІНТЕГРАЦІЯ МЕТОДІВ ШТУЧНОГО ІНТЕЛЕКТУ В РОЗРОБКУ ОРІЄНТОВАНИХ НА КОРИСТУВАЧА ІНТЕРФЕЙСІВ

Т.А. Lypak, post-graduate student

Ternopil Ivan Pul'uj National Technical University, Ukraine

INTEGRATION OF ARTIFICIAL INTELLIGENCE METHODS INTO THE DESIGN OF USER-CENTERED INTERFACES

У сучасному цифровому середовищі інтерфейси користувача є не лише точкою входу до сервісів, але й інструментом формування лояльності, довіри та ефективності взаємодії користувача із системою. Інтеграція методів штучного інтелекту (ШІ) у проєктування орієнтованих на користувача інтерфейсів дозволяє створювати інтерактивні, адаптивні та персоналізовані цифрові середовища. Проведене дослідження мало на меті проаналізувати можливості та практики інтеграції ШІ в процеси проєктування інтерфейсів, що орієнтовані на поведінку, потреби та контекст користувача. Для цього було здійснено аналіз літератури з UX-дизайну та штучного інтелекту; досліджено приклади застосування ШІ у UI/UX-дизайні; проведено порівняльний аналіз моделей персоналізації інтерфейсів та узагальнено перспективні напрями застосування ШІ в розробці інтерфейсів.

Одним з важливих напрямів інтеграції штучного інтелекту в UI/UX є персоналізація інтерфейсу. Успішно застосовуються рекомендаційні системи на основі колаборативної фільтрації та глибокого навчання [1]. Для прогнозування поведінки користувачів, зокрема, послідовності їх дій, використовуються довга короткочасна пам'ять та рекурентні нейронні мережі [2]. Стрімко набирають популярності голосові й текстові інтерфейси, для розробки яких застосовують великі мовні моделі, наприклад, ChatGPT, DialogFlow, Alexa API тощо [3]. Алгоритми комп'ютерного зору підтвердили свою ефективність у UX-аудиті, для аналізу емоцій і міміки користувачів [4]. Проведення UX-аналітики в реальному часі значно полегшують теплові карти на базі штучного інтелекту, поведінкові дашборди, A/B тестування з нейромережевими прогнозами [5].

Наведемо приклад застосування ШІ у смарт-застосунку для міського транспорту: ШІ моделює поведінку користувача на основі історії поїздок, погоди, часу; інтерфейс адаптується в реальному часі - показує релевантні маршрути, змінює розміщення кнопок.

Інтеграція ШІ в UX-дизайн інтерфейсів забезпечує динамічну персоналізацію, покращує контекстну взаємодію; створює передумови для когнітивно адаптивних систем у цифровому середовищі. Однак таке повсюдне впровадження рішень ШІ породжує нові виклики: етичність, прозорість рішень, контроль за автономними адаптаціями.

Література

1. García-González A. et al. User experience and AI-based personalization. Computers in Human Behavior, 2022.
2. Li C. et al. User Behavior Prediction with Recurrent Neural Networks. ACM Transactions, 2021.
3. OpenAI. ChatGPT Technical Report. 2023. <https://openai.com/research>

4. Valstar M. et al. Affect recognition in HCI. In Human-Centered AI. Springer, 2020.
5. Huang J. et al. Real-time AI-enhanced UX analytics. IJHCS, 2023.

УДК 004.8

Г.І. Липак, к.н.с.к., доцент, Д.А. Коломийчук

Тернопільський національний технічний університет імені Івана Пулюя

ВПЛИВ БЕКДОР АТАК НА МОЖЛИВОСТІ НАВЧАННЯ НЕЙРОННИХ МЕРЕЖ

H. Lypak, Ph.D., D. Kolomyichuk

Ternopil Ivan Puluj National Technical University

THE IMPACT OF BACKDOR ATTACKS ON THE LEARNING CAPABILITY OF NEURAL NETWORKS

У роботі було проаналізовано непомітні атаки зловмисників на дані часових рядів у рекурентних нейронних мережах, щоб вивчити як безпеку глибоких рекурентних нейронних мереж, так і зрозуміти властивості навчання в глибоких рекурентних нейронних мережах.

Оскільки глибокі нейронні мережі широко використовуються в прикладних областях, існує можливість зниження точності та безпеки за допомогою методів зловмисників. Метод зловмисників - це отруєння даних через бекдор, коли зловмисник збиває навчальні вибірки невеликим збуренням, щоб неправильно класифікувати вихідний клас до цільового класу. При зміні даних через бекдор зловмисник має доступ до підмножини навчальних даних з мітками, можливість впливати на навчальні вибірки та можливість змінити мітку вихідного класу на мітку цільового класу. Зловмисник не має доступу до класифікатора під час навчання або знання процесу навчання. Також було досліджено захист від зміни даних через бекдор після навчання, розглядаючи ітераційний метод для визначення пари вихідного та цільового класів у такій атаці.

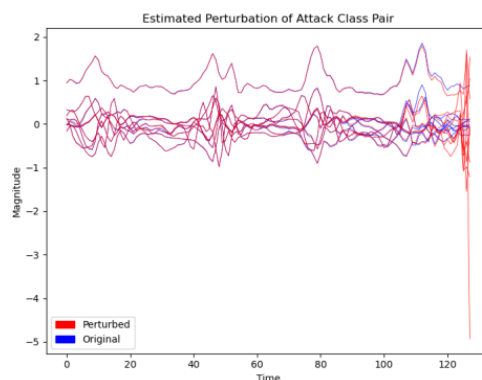


Рисунок 1 - Очікуване збурення пари класів атаки, LSTM

Методи атак через бекдор, успішно обманюють класифікатор LSTM, не знижуючи точності тестових зразків, без наявності шаблону бекдору. По-друге, метод захисту успішно визначає пару вихідних класів в такій атаці. По-третє, отруєння через бекдор у LSTM вимагає або більше навчальних зразків, або більшого збурення, ніж стандартна мережа прямого зв'язку.

Література.

2. B. Tran, J. Li and A. Madry. Spectral Signatures in Backdoor Attacks // NeurIPS. 2018. URL: <https://arxiv.org/abs/1811.00636>.