

ПРОБЛЕМИ ЗАБЕЗПЕЧЕННЯ БЕЗПЕКИ ЖИТТЄДІЯЛЬНОСТІ ЛЮДИНИ, ПОВ'ЯЗАНІ З РОЗВИТКОМ ШТУЧНОГО ІНТЕЛЕКТУ

*Проскученко Р.С., студент (гр. РІ-41, РТФ КПІ ім. Ігоря Сікорського);
Гусєв А.М., к.б.н., доцент (каф. ОПЦБ КПІ ім. Ігоря Сікорського)*

Вступ. В ХХІ ст. в інформаційну епоху, супроводжувану комп'ютерною та цифровою революцією, зміною технічної реальності в ході четвертої промислової революції, людина стикається зі збільшенням загроз життєдіяльності людства: різке зростання ймовірності природних, соціальних, техногенних і космічних катастроф. Нагальним завданням є інформаційна безпека, особливо у зв'язку з процесами симбіозу людини і техніки, в ході роботизації, кіборгізації, впровадження об'єктів з штучним інтелектом.

Позитивними прикладами систем із штучним інтелектом, які застосовуються в даний час, можуть служити системи автономного управління. Широко використовуються медичні діагностичні програми. Відомі успіхи в робототехніці, кіборгізації, в розумінні елементів природної мови. Блискуче показали себе об'єкти з штучним інтелектом у військовій сфері, в системах протиповітряної оборони, у проведенні кібервійн, а також для забезпечення інших завдань національної безпеки.

Предметом дослідження є загальна оцінка ризиків та проблема створення механізмів контролю штучного інтелекту. Розгляд вищезазначених питань буде проводитись з урахуванням вимог охорони праці.

Аналіз публікацій. Штучний інтелект, імітуючи розумову активність людини, досяг досить високого рівня і відмінно виконує завдання з розпізнавання, класифікації, прийняття рішень і т. д. Подібні завдання вирішуються і при забезпеченні безпеки, а так як штучний інтелект дозволяє це робити точніше і швидше, то його широко використовують для побудови систем безпеки життєдіяльності людини. До таких систем можна віднести сенсорні мережі, «розумну» техніку, «розумні» будинки. Можна вважати, що в цьому напрямку розробка штучного інтелекту в цілях оптимізації та забезпечення безпеки життєдіяльності людини здійснюється успішно [1, с. 108].

Основні результати дослідження. Наскільки широко людство може використовувати досягнення в області розвитку штучного інтелекту, не нашкодивши своїй біологічній та соціокультурній сутності? Експеримент вчених з університету Кента (Великобританії) виявив вплив GPS-навігаторів на розумові процеси людини [2]. Вчені припустили, що, позбавляючи людину від будь-яких інтелектуальних функцій за рахунок впровадження системи штучного інтелекту, ми «розвантажуюмо» мозок від рутинних дій, але «відпочиваючі» частини мозку, які відповідали за цю діяльність, можуть почати деградувати.

Зворотною стороною передачі безпеки життєдіяльності людини штучному інтелекту є той факт, що людина потребує забезпечення безпеки постійно, тому і відповідні системи повинні працювати безперервно. Однак подібні системи дуже складні, їх алгоритми чутливі до найменших змін. А чим

складніше система, тим більше чинників, які можуть вивести її з нормального функціонування. Програмне забезпечення систем буде постійно застарівати, технічний пристрій – старіти, виходити з ладу.

Якщо побудувати систему забезпечення безпеки життєдіяльності людини тільки на використанні технічних засобів, то завжди буде ймовірність виходу з ладу подібної техніки, що веде до високої вразливості захищених об'єктів. В середині травня 2017 р. світ потрясло зараження мережним вірусом WannaCrypt. З його допомогою виявилися заражені і виведені з ладу понад 200 тисяч комп'ютерів в десятках країн світу (причому очевидно, що всі офіційні державні сайти та сайти провідних корпорацій володіють стійкими системами захисту інформації у власних мережах) [3].

Останнім часом набуває популярності такий феномен, як «розумний» будинок [4]. Він являє собою автоматизовану інтелектуальну систему, спрямовану на підвищення якості життя людини за допомогою високотехнологічних пристроїв. Наслідки від «хакерських атак» на «розумний» будинок можуть бути непередбачуваними, наприклад, зловмисники отримують доступ до важливих елементів будинку, таким, як сервери зберігання інформації, засоби комунікації, управління електрикою, опаленням, вентиляція і т. д. В результаті «розумний» будинок з комфортного і безпечного житла в одну мить може перетворитися в «пастку», що несе загрозу життя мешканців котеджу. Повністю покладатися на інформаційні технології, об'єкти з штучним інтелектом у таких проектах не можна, повинна існувати паралельна система «ручного захисту» і можливість перемикання на запасну, постійно поновлювану (електронну) систему. Таким чином, процес захисту «розумних» будинків стає пролонгованим, неможливо створити ідеальну інтелектуальну системи захисту, коли проти неї «грає» творчий розум людини.

Відомий вчений С. Хокінг, аналізуючи поведінку сучасної людини і використання передових технологій, також висловив побоювання, що широке застосування штучного інтелекту може перетворитися за своїми наслідками для людей, які живуть в інформаційному суспільстві, в найгірше, що коли-небудь могло статися з людством. Він неодноразово говорив, що всі переваги попередніх винаходів незрівнянні з тим, що може статися найближчим часом у процесі використання об'єктів із штучним інтелектом. С. Хокінг вважає, що розвиток програмного забезпечення, орієнтований на самовдосконалення штучного інтелекту, може перетворитися на катастрофу для людства. А якщо при цьому штучний інтелект буде контролювати безпеку життєдіяльності людини, то людство просто виявиться беззахисним в цьому протистоянні [5].

На початку січня 2017 р. в Каліфорнії проходила конференція Beneficial AI 2017. Результатом конференції стали «Азіломарські принципи штучного інтелекту», в рамках яких повинні відбуватися подальші дослідження у сфері штучного інтелекту.

Основні положення цих принципів (скорочено):

- дослідження в сфері штучного інтелекту має бути створення не некерованого, а корисного інтелекту безпечного і надійного протягом усього терміну експлуатації;

- розробники штучного інтелекту грають ключову роль у формуванні моральних наслідків його і несуть обов'язок впливати на такі наслідки;

- участь автономної системи в процесі прийняття судових рішень повинна супроводжуватися наданням переконливих пояснень, які можуть бути перевірені ще раз людьми з компетентних органів влади;

- системи штучного інтелекту повинні розроблятися і працювати таким чином, щоб бути сумісними з ідеалами людської гідності, його правами і свободами;

- системи штучного інтелекту, що розробляються з можливістю рекурсивного самовдосконалення або самовідтворення з наступним швидким збільшенням їх кількості або якості, повинні відповідати суворим критеріям безпеки і контролю

- суперінтелект повинен розроблятися тільки для служіння на благо всього людства, а не однієї держави або організації [6].

Оскільки управління системами з штучним інтелектом вимагають підготовки висококваліфікованих кадрів, не варто забувати, що фахівець повинен володіти не тільки професійними знаннями, він повинен вміти мислити, бути здатним змінювати стереотипи сприйняття дійсності, а, отже, – володіти глибокою гуманітарною підготовкою. Доводиться визнати, що якісна освіта стає чинником безпеки життєдіяльності людства [7].

Висновки. До процесу передачі забезпечення безпеки життєдіяльності людини в руки штучного інтелекту потрібно ставитися дуже обережно і постійно контролювати всі можливі наслідки. Можна рекомендувати фахівцям використовувати різні пристрої, що різняться за принципом дії, щоб вони могли доповнювати один одне і стабілізувати систему в цілому. Слід створювати програми, що виключають можливість нанесення пристроєм з штучним інтелектом шкоди людині і його безпеці. Найважливішим елементом системи безпеки життєдіяльності повинна залишатися сама людина, її природний розум, розвинений інтелект, здатність до творчості. Сам же штучний інтелект як елемент нашого життя постійно створює перед окремою людиною, державою, суспільством в цілому ситуацію перманентного вибору принципово нових варіантів, які були відсутні в перехідних станах попередніх епох [8].

Література

1. Абдулатипова М.А., Камалова Р.Ш. Искусственный интеллект // Международный научно-исследовательский журнал. 2013. №. 5 (12). С. 108.

2. Roger McKinley. Technology: Use or lose our navigation skills // Nature. 2016. № 531. P.573–575.

3. Официальный сайт Microsoft Corporation: The need for urgent collective action to keep people safe online: Lessons from last week's cyberattack. Режим

доступу: <https://blogs.microsoft.com/on-the-issues/2017/05/14/need-urgent-collective-action-keep-people-safe-online-lessons-last-weeks-cyberattack/> (дата звернення: 25.10.2017).

4. Нехамкин В.А., Сачкова В.А. Образы «города будущего» в культуре конца XX – начала XXI вв.: мифы и реальность Режим доступа: <https://elibrary.ru/item.asp?id=18280789> (дата звернення 25.10.2017).

5. Официальный сайт британского издания «Daily Express». Stephen Hawking: AI could be the worst thing that ever happened to humanity. Режим доступа: <http://www.express.co.uk/news/science/723072/Stephen-Hawking-AI-end-of-the-worldworst-thing-humanity/> (дата звернення 25.10.2017).

6. Азиломарские принципы искусственного интеллекта. Режим доступа: <http://robotrends.ru/pub/1737/azilomarskie-principy-iskusstvennogo-intellekta> (дата звернення 25.10.2017).

7. ADVANCED EDUCATION Режим доступа: <http://ae.fl.kpi.ua/> (дата звернення 25.10.2017).

8. Нехамкин В.А. Проблема исторического выбора: личностно-психологический аспект // European Social Science Journal. 2015. № 6. С. 390-397.