

ВСТУП

Поняття “*машинне навчання*” як науковий напрям виникло при зародженні *штучного інтелекту* (рис. В1), коли деякі дослідники були зацікавлені в навчанні машин на даних з певної бази [В1]. Були спроби підійти до проблеми за допомогою різних *символічних методів*, а також тогочасних “нейронних мереж”, якими були здебільшого *персептрони* та інші моделі, які пізніше виявились наново відкритими *узагальненими лінійними моделями статистики* [В2], а також використання ймовірнісних підходів, особливо в автоматизованій медичній діагностиці [В3]: Однак посилення акценту на логічному, основаному на знаннях, підході спричинило розрив між системами *штучного інтелекту* та *машинним навчанням*.

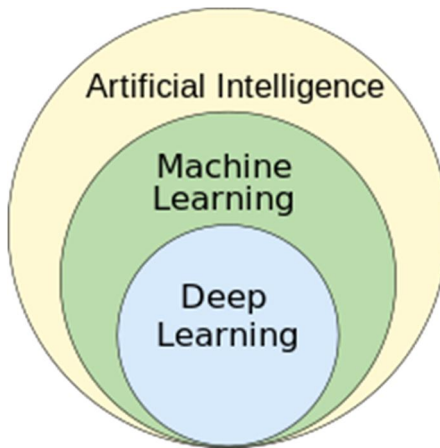


Рис. В1. Співвідношення понять “штучний інтелект”, “машинне” та “поглиблене навчання”

До 1980 р. *експертні системи* почали домінувати над штучним інтелектом [В4]; робота над символічним навчанням, базованим на знаннях, продовжувалася в межах штучного інте-

лекту, що зумовило виникнення *індуктивного логічного програмування* (*Inductive Logic Programming – ILP*), але суто статистична лінія досліджень тепер була поза межами власне штучного інтелекту, у сфері розпізнавання образів та пошуку інформації [B3].

Штучний інтелект та інформатика припинили дослідження нейронних мереж приблизно одночасно. Цю лінію продовжували за межами області штучний інтелект / комп'ютерні науки (*Artificial Intelligence / Computer Science – AI/CS*), як “конекціонізм”, дослідники з інших дисциплін, зокрема, Гопфілд (*Hopfield*), Румельхарт (*D. E. Rumelhart*) та Гінтон (*G. E. Hinton*). Їхній головний успіх настав у середині 1980-х років із повторним винаходом зворотного поширення похибки [B3].

Машинне навчання – це *підгалузь інформатики*, яка розвинулася з вивчення теорії розпізнавання образів і обчислювального навчання в області штучного інтелекту. У 1959 р. Артур Самуель (*Arthur Samuel*) визначив машинне навчання як “область дослідження, яка дає комп'ютерам можливість навчатися без явного програмування”. Машинне навчання, реорганізоване та визнане окремою сферою, почало процвітати в 1990-х роках. Сфера змінила свою мету з досягнення штучного інтелекту на вирішення практичних проблем, фокус змістився від символічних підходів, успадкованих від штучного інтелекту, до методів і моделей, запозичених із статистики, нечіткої логіки та теорії ймовірностей [B4]. Хоча найперша модель машинного навчання була представлена в 1950-х роках, коли Артур Самуель написав програму, яка обчислювала шанси на перемогу в шашках для кожної сторони, історія машинного навчання сягає корінням у десятиліття людського бажання та зусиль щодо вивчення когнітивних процесів людини [B5]. Зокрема, у 1949 р. канадський психолог Дональд Гебб (*Donald Hebb*) опублікував книгу “Організація поведінки”, в якій представив теоретичну нейронну структуру, утворену певними взаємодіями між нервовими клітинами [B6]. Модель Гебба про взаємодію нейронів

один з одним заклала основу пояснення того, як штучний інтелект і алгоритми машинного навчання працюють у вузлах або штучних нейронах, які використовуються комп'ютерами для передачі даних [B5]. Інші дослідники, які вивчали людські когнітивні системи, також зробили внесок у сучасні технології машинного навчання, зокрема логік Уолтер Піттс (*Walter Pitts*) та Воррен Мак-Каллох (*Warren McCulloch*), які запропонували перші математичні моделі нейронних мереж, щоб створити алгоритми, які відображають процеси мислення людини [B5].

На початку 1960-х років компанія *Raytheon* розробила експериментальну “навчальну машину” з пам'яттю на перфострічці під назвою “Кібертрон” для аналізу сигналів гідролокатора, електрокардіограм і мовленнєвих моделей за допомогою елементарного навчання з підкріпленням. Людина-оператор / вчитель постійно “навчала” її розпізнавати закономірності та користувалася кнопкою повторного змушування її оцінювати неправильні рішення [B7]. Репрезентативною книгою про дослідження машинного навчання в 1960-х роках була книга Нільссона (*Nilsson*) про навчальні машини (*Learning Machines*), присвячена переважно машинному навчанню для класифікації шаблонів [B8]. Інтерес до розпізнавання образів тривав у 1970-ті роки, як описали Дуда (*Duda*) та Гарт (*Hart*) у 1973 р. [B9]. У 1981 р. було виголошено доповідь про використання стратегій навчання, які би дозволили штучній нейронній мережі навчитися розпізнавати 40 символів (26 літер, 10 цифр і 4 спеціальні символи) з комп'ютерного терміналу [B10]. Том М. Мітчелл (*Tom M. Mitchell*) надав формальне визначення алгоритмів, що вивчаються в області машинного навчання [B11], яке відповідає пропозиції Алана Тюрінга (*Alan Turing*) замінити запитання “Чи можуть машини мислити?” запитанням “Чи можуть машини робити те, що ми (як мислячі сутності) можемо робити?” [B12] в його статті “Обчислювальна техніка та інтелект”.

Сучасне машинне навчання має дві мети. Одна полягає в класифікації даних на основі розроблених моделей; інша – в тому, щоб спрогнозувати майбутні результати на основі цих

моделей. Гіпотетичний алгоритм для класифікації даних може використовувати комп'ютерний зір у поєднанні із навчанням під наглядом, щоб навчити його класифікувати, наприклад, ракові плямки. Алгоритм машинного навчання для біржової торгівлі може інформувати трейдера про майбутні потенційні прогнози.

Конекціоністський підхід до побудови систем штучного інтелекту розвинувся на противагу символічному, характерного для сучасних *моделей знань*. В основі конекціоністського підходу – спроба безпосереднього моделювання розумової діяльності людського мозку. Відомо, що мозок людини складається з величезної кількості нервових клітин (нейронів), що взаємодіють між собою. Ці “обчислювальні елементи” мозку функціонують набагато повільніше, ніж обчислювальні елементи комп'ютерних систем. Але ефективність людського інтелекту досягається як завдяки паралельній роботі нейронів, так і тому, що механізми їх взаємодії були вироблені протягом тривалої еволюції.

Штучні нейронні мережі, або конекціоністські системи, є обчислювальними системами, що намагалися наслідувати біологічні нейронні мережі людського мозку, проте й відрізняються від нього. Зокрема, штучні нейронні мережі мають тенденцію бути *статичними* та *символічними*, тоді як біологічний мозок більшості живих організмів є *динамічним (пластичним)* та *аналоговим* [B13]. Машинне навчання та розпізнавання образів можна розглядати як дві грані однієї сфери. Штучні нейронні мережі, як правило, розглядають як низькоякісні моделі функції мозку. Такі системи “вчаться” моделювати складні функції, наприклад, для класифікації, трансляції, відповідей на запитання, прогнозування, на основі достатньо великого набору навчальних даних. Мережа складається із вузлів (вершин), званих *нейронами*, які з'єднані з іншими нейронами *ребрами*, які за функціями відповідають *синапсам* у мозку. Нейрони продукують вихідний сигнал, представлений (дійсним) числом. Значення цього числа обчислюється за *функцією активації*, аргументами якої є значення вихідних сигналів, що передаються від

інших нейронів через вхідні порти нейрона. Вихід кожного штучного нейрона обчислюється за допомогою деякої нелінійної функції від суми його входів. Штучні нейрони та ребра зазвичай мають *ваги*, значення яких визначається даними навчання в процесі навчання. Вага збільшує або зменшує силу сигналу на певному з'єднанні (ребрі). Штучні нейрони можуть мати певний *пори́г*, тобто сигнал надсилається на вихід нейрона лише тоді, коли сукупний вхідний сигнал перетинає цей поріг. Як правило, штучні нейрони розміщені шарами. Початкова мета створення штучних нейронних мереж полягала в тому, щоб вирішувати проблеми так само, як це робив би людський мозок. Однак з часом увага перемістилася на виконання конкретних завдань, що призвело до відхилень від суто біологічних принципів.

Багатошарові штучні нейронні мережі називають також *глибокими (поглибленими)* нейронними мережами. Різні шари можуть виконувати різні типи перетворень на своїх входах. Сигнали проходять від першого шару (*вхідний рівень*) до останнього (*вихідний рівень*), можливо, після кількох проміжних шарів. Багатошарові мережі з алгоритмами глибокого навчання стали найуспішнішим інструментом машинного навчання.

Поглиблене навчання забезпечується кількома прихованими шарами штучної нейронної мережі. Цей підхід намагається змоделювати те, як людський мозок перетворює світло та звук на зір та слух. Поглиблене навчання – це назва класу *алгоритмів машинного навчання*, які використовують кілька шарів для поступового вилучення вищих шарів функцій на різних рівнях абстракції з неопрацьованих вхідних даних. Наприклад, під час опрацювання мови нижчі рівні (шари) можуть ідентифікувати типи слів, тоді як вищі рівні ідентифікують граматичні зв'язки; під час опрацювання зображень нижчі рівні можуть ідентифікувати краї, тоді як вищі – поняття, пов'язані з розпізнаванням цифр, літер чи обличчя. Найглибше навчання належить до повністю автоматичного

навчання від джерела до остаточно вивченого об'єкта. У програмі розпізнавання зображень неопрацьований вхід може бути матрицею пікселів; перший репрезентативний шар може абстрагувати пікселі та кодувати краї; другий рівень може складати та кодувати розташування країв; третій шар може кодувати ніс і очі; а четвертий рівень – розпізнати, що зображення містить обличчя. Важливо, що процес поглибленого навчання може самостійно дізнатися, які функції на якому рівні оптимально розмістити. Це не усуває необхідності ручного налаштування; наприклад, різна кількість шарів і розміри шарів можуть забезпечити різний ступінь абстрагування [B14]. Під час поглибленого навчання кожен рівень вчиться перетворювати свої вхідні дані на дещо абстрактніше та складніше подання.

Термін упровадила в 1986 р. Ріна Дехтер (*Rina Dechter*) для машинного навчання, відтак у 2000 р. Ігор Айзенберг (*Igor Aizenberg*) та його колеги для штучних нейронних мереж. Сьогодні поглиблене навчання майже виключно належить до машинного навчання на основі *глибоких* штучних нейронних мереж, тобто багатшарових штучних нейронних мереж, в яких шари є *векторами вершин*, так званих нейронів. Поглиблене навчання може бути контрольованим, напівконтрольованим або неконтрольованим. Під час процесу навчання вагові коефіцієнти періодично коригуються відповідно до зміни (варіацій) даних. Зазвичай вагові коефіцієнти коригуються за допомогою зворотного поширення похибок через шари у зворотній послідовності. Для покращення роботи мережі вагові коефіцієнти оптимізують за допомогою *методу градієнтного спуску* доти, доки мережа не почне виконувати завдання із задовільною похибкою (точністю). Серед успішних застосувань поглибленого навчання – комп'ютерний зір та розпізнавання мови

Глибокі штучні нейронні мережі, як правило, є мережами прямого зв'язку, в яких дані переходять із вхідного рівня на вихідний без повернення назад.

Перший загальний, робочий алгоритм навчання для керованих, глибинних, прямих, багат шарових перцептронів опублікували Олексій Івахненко і Валентин Лапа в 1967 р. [B15] Перший багат шаровий перцептрон поглибленого навчання, навчений стохастичним градієнтним спуском [B16], опублікував у 1967 р. Шунічі Амарі (*Shun'ichi Amari*) [B17], [B18]. Архітектури поглибленого навчання для згорткових нейронних мереж (CNN) зі згортковими шарами та рівнями зменшення дискретизації почалися з неокогнітрона (*Neocognitron*), який запропонував Куніхіко Фукусіма (*Kunihiko Fukushima*) в 1980-х роках [B19].

У поглибленому навчанні застосовують різні типи нейронних мереж:

- Перцептрон (*perseptron*) – найпростіша форма, яка використовується для бінарної класифікації.
- Згорткові нейронні мережі (*CNN – Convolutional Neural Networks*) – ефективні в опрацюванні зображень, виділяють ознаки, як-от грані та текстури.
- Рекурентні нейронні мережі (*RNN – Recurrent Neural Networks*) – надаються для послідовних даних, із пам'яттю для урахування попередніх станів.
- Глибокі нейронні мережі (*DNN – Deep Neural Networks*) – багат шарові мережі, що виявляють складні образи (патерни) в даних.
- Спеціалізовані мережі (наприклад, *GAN* і *LSTM*) – генеративно-змагальні мережі (*GAN – Generative adversarial network*) для створення даних і довга короткочасна пам'ять (*LSTM – Long short-term memory*) для опрацювання послідовних даних із довгостроковою залежністю.

Різноманітність цих мереж дає змогу поглибленому навчанню успішно розв'язувати такі задачі, як розпізнавання образів і опрацювання тексту.

Розглянемо докладніше відмінності глибинного навчання (*Deep Learning*) від традиційного машинного навчання [B20–B22] (див. таблицю).

**Відмінності поглибленого навчання (*Deep Learning*)
від традиційного машинного навчання**

Характеристика	Deep Learning	Машинне навчання
1	2	3
Архітектура моделей	Глибокі нейронні мережі з безліччю шарів, здатних видобувати складні ознаки	Поверхневі моделі з невеликою кількістю шарів, найчастіше лінійні або дерево-подібні
Використання даних	Потребує великих обсягів даних для ефективного навчання	Може працювати з меншими обсягами даних, часто ефективніше в завданнях з обмеженими даними
Робота з ознаками	Автоматичне вилучення ознак у процесі навчання	Потребує ручної інженерії ознак для визначення найкращих характеристик
Навчання	Ієрархічне, з уточненням уявлень на різних рівнях	Зазвичай плоскіше, без ієрархії в поданні даних
Опрацювання даних	Ефективно опрацьовує неструктуровані дані, такі як зображення, тексти, звук	Більш орієнтований на структуровані дані, такі як таблиці та бази даних
Продуктивність	Потребує потужних обчислювальних ресурсів, особливо блок опрацювання графіки (<i>GPU – graphics processing unit</i>)	Може працювати на менш потужних пристроях, що робить його доступнішим

Продовження табл.

1	2	3
Інтерпретова- ність	Моделі можуть бути складними і важкозрозумілими ("чорна скринька")	Простіші моделі, легше інтерпретувати і пояснити
Застосування	Ефективний у завда- ннях, що потребують опрацювання великих обсягів даних і виді- лення складних закономірностей	Часто використо- вується в задачах з обмеженими даними і явними ознаками

Ці відмінності відображають основні особливості обох підходів, надаючи змогу вибрати відповідний залежно від вимог завдання і доступних даних.