

ОБРАЗНИЙ ПІДХІД ДО ОБЧИСЛЕННЯ КІЛЬКОСТІ ІНФОРМАЦІЇ ТА ОЦІНКИ ЇЇ ЦІННОСТІ

© Заяць В. М., Заяць М. М., 2017

Розглянуто різні підходи до визначення та обчислення основних понять теорії інформації, зокрема кількості інформації та оцінки її цінності на підставі статистичних міркувань (класичний підхід), теорії алгоритмів (алгоритмічний підхід) та теорії розпізнавання образів (образний підхід). Здійснено їх порівняльний аналіз та відзначено сфери їх ефективного застосування. Запропоновано новий підхід, який дає змогу кількісно оцінити цінність інформації.

Ключові слова: теорія інформації, кількість інформації, ймовірність, алгоритм, образ, цінність інформації.

Existent approaches to determination of basic concepts of information theory, coming from the statistical reasoning (classic approach), theory of algorithms (algorithmic approach) and theory of pattern recognition (vivid approach) are examined. Their comparative analysis is conducted. Approach which enables in number to estimate the value of information is offered.

Key words: theory of information, quantity of information, probability, algorithm, pattern, value of information.

Вступ. Постановка проблеми

Сьогодні сформовано цілу множину означень терміна “інформація” (від лат. *informatio* – роз’яснення, подання). Від найзагальнішого – відображення реального світу до найвужчого – сукупність даних, повідомлень, відомостей, знаків, образів, які підлягають опрацюванню, впорядкуванню, та отримання нових знань і їх використання. Оскільки це поняття вживається практично у всіх сферах людської діяльності, то залежно від цілей дослідження використовують те або інше означення. З практичного погляду найдоцільнішим видається означення інформації французького вченого Буазе, як всього того, що знімає невизначеність з наших знань про даний об’єкт, явище або процес. Отже, інформація є сирим матеріалом, який полягає у збиранні даних. Формування ж знань передбачає певні розмірковування, які впорядковують отримані дані на основі їх зіставлення та класифікації. Для формування наукової теорії потрібно дати чіткі й однозначні визначення, які виключають будь-які двозначності в межах того розділу науки, в якому їх застосовують.

Точні означення термінів у науковій мові ґрунтуються на двох підходах. Перший підхід (теоретичний) часто застосовується у математичних дисциплінах, де визначення розпочинається із формулювання основних постулатів або гіпотез. Всі інші складніші сутності виводяться із цих початково введених тверджень і їм не суперечать. Другий підхід до формулювання означень характерний для експериментальних наук і часто називається операційним (практичним). За такого підходу вважається доцільним вводити у наукову мову ті величини, які можна визначити експериментально або оцінити опосередковано через їх вплив на поведінку досліджуваної системи. Нововведені терміни, які не піддаються операційному визначенню і не мають практичної цінності, вилучаються з наукового словника. Зокрема, так було з поняттям ефіру, якого теорія відносності позбавила будь-якого змісту і він вийшов з ужитку в науковій мові, незважаючи на численні публікації з цієї проблеми.

Формування цілей

Мета роботи – порівняти відомі методи та підходи до визначення кількості інформації та оцінки її цінності, вказати межі їх застосувань та запропонувати новий підхід до кількісної оцінки цінності інформації.

Виклад основного матеріалу

У теорії інформації основним поняттям є поняття кількості інформації як функції відношення можливих станів системи до отримання і після отримання інформації про неї, взятої у логарифмічному масштабі з метою забезпечення властивості адитивності інформації. Незважаючи на появу нових визначень цього поняття, відповідно до теорії алгоритмів та методів теорії розпізнавання образів, які дають змогу розв'язати деякі часткові прикладні задачі, класичне визначення залишається актуальним і методи цієї теорії з успіхом застосовуються до проблем кодування інформації, зв'язку, нечіткої логіки, комп'ютерної техніки. Ця теорія корисна з погляду формулювання правил і чіткого визначення, що в її межах можна, а чого не можна зробити. Зокрема, неможливо в межах цієї класичної теорії інформації дослідити процес мислення, оскільки не розглядається оцінка інформації людиною.

Будь-яка фізична система не є повністю визначеною. Ми маємо дані про деякі макроскопічні змінні системи, але не можемо точно визначити стани і швидкості всіх атомів, що утворюють систему. Мірою невизначеності станів системи є поняття ентропії. Ентропія виражає кількість відсутньої інформації про ультрамікроскопічну структуру системи. Ці два поняття кількості інформації та ентропії слід розглядати і трактувати сумісно. Ця думка сформульована як негентропійний принцип інформації у роботі [1]. Суть цього принципу зводиться до твердження: отримання інформації про систему пов'язане зі зменшенням її ентропії. Наскільки зменшиться ентропія системи, настільки зросте інформація про неї. З цієї позиції інформацію можна означити як негативну ентропію системи.

Класичний підхід

У класичній теорії інформації, яка розроблена в фундаментальних роботах Р. Хартлі [2] та К. Шенона [3], основні поняття ґрунтуються на теорії ймовірностей. Згідно з теорією Р. Хартлі, якщо система може перебувати в N станах, то повна кількість інформації про неї визначається як двійковий логарифм від кількості цих станів:

$$I = \log_2 N. \quad (1)$$

Одиницею вимірювання кількості інформації є один біт. Отже, біт інформації можна визначити як кількість інформації, що міститься у повідомленні з двома можливими станами. Зауважимо, що ймовірності появи цих повідомлень можуть бути різними. Якщо априорі невідома кількість можливих станів системи, але кожен із них формується за допомогою n різноякісних ознак, кожна з яких може бути на одному з k розрядів, то, скориставшись формулою розміщень C_k^n та застосувавши формулу Стірлінга для обчислення логарифма від факторіала

$$\ln X! \approx X(\ln X - 1),$$

яка є справедливою для великих значень X , отримуємо формулу Шенона [3] для визначення кількості інформації:

$$I = -k \cdot \sum_{i=1}^n p_i \cdot \log_2 p_i, \quad (2)$$

де p_i – ймовірність появи i -ї якісної ознаки. Якщо поява кожної із ознак є рівноімовірною, то

$$p_i = \frac{1}{n} \text{ і з (2) одержимо}$$

$$I = -k \cdot \log_2 \frac{1}{n} = k \cdot \log_2 n, \quad (3)$$

що відповідає формулі (1), оскільки $N = n^k$.

Отже, формули (1), (2) і (3) дають змогу вираховувати кількість інформації як за рівномірної, так і за нерівномірної ймовірності появи якісних ознак незалежно від якості інформації, що отримана (повідомлення, звук, зображення чи сигнал будь-якої природи). Але при цьому не враховується цінність інформації, яка є значною мірою суб'єктивною характеристикою і залежить від цілей та уподобань користувача. По суті, це та плата за можливість розрахунку кількості інформації, на підставі статистичних міркувань, незалежно від її якості та цінності.

Зазначимо, що кількість корисної інформації є завжди величиною додатною. У випадку взаємного впливу однієї системи на іншу з'являється умовна невизначеність станів системи, яка

може привести до зміни знака у визначенні кількості інформації. Тоді маємо справу з іншою якістю інформації – дезінформацією, або хибною інформацією.

На перший погляд, надуманість визначення (1) за допомогою логарифмічної функції цілком виправдана з прикладної позиції, оскільки забезпечує властивість адитивності інформації. Справді, нехай маємо Z підсистем, що взаємодіють, кожна з яких може перебувати в N_j станах, де $j = 1, 2, \dots, Z$, об'єднаних в одну систему. Якщо стани однієї підсистеми ніяк не впливають на стани іншої підсистеми, тобто є взаємно незалежними, то кількість всіх можливих станів, у яких може перебувати сукупність аналізованих систем, визначаємо як

$$N = N_1 \cdot N_2 \cdot N_3 \cdot \dots \cdot N_Z,$$

а з урахуванням (1) кількість отриманої інформації про систему

$$I = \text{Log}_2(N_1 \cdot N_2 \cdot N_3 \cdot \dots \cdot N_Z) = \text{Log}_2 N_1 + \text{Log}_2 N_2 + \dots + \text{Log}_2 N_Z = I_1 + I_2 + \dots + I_Z$$

Отже, кількість інформації, яку можна отримати від довільної кількості взаємонезалежних систем, що взаємодіють, визначається сумою кількостей інформації, які можна отримати про кожен з них зокрема.

Аналізуючи вираз (2), можна встановити умови, за яких кількість інформації набуває екстремальних значень. Складові під знаком суми у виразі (2) набувають мінімального (нульового) значення за двох випадків:

- а) ймовірність появи p_i ознаки дорівнює одиниці, що впливає з виразу (2);
- б) ймовірність появи p_i ознаки дорівнює нулю, в чому можна переконатися, застосувавши правило Лопітала для розкриття невизначеності:

$$\lim_{p_i \rightarrow 0} (-p_i \cdot \text{Log}_2 p_i) = \lim_{k \rightarrow \infty} \frac{\text{Log}_2 k}{k} = \lim_{k \rightarrow \infty} \frac{\text{Log}_2 e}{k} = 0.$$

Скориставшись методом неозначених множників Лагранжа стосовно функції, що стоїть під знаком суми у виразі (2), можна переконатися, що кількість інформації досягає максимального значення у випадку, якщо поява всіх якісних ознак є рівноімовірною. Отже, формула (1) дає можливість визначити максимальну кількість інформації, якщо кожен із N станів системи має однакову ймовірність появи.

За наявності взаємного зв'язку між системами A і B , що взаємодіють, вводиться поняття умовної ентропії, яка визначає невизначеність появи деякого стану системи A за умови, що перед тим з'явився деякий стан системи B :

$$H(A/B) = \sum_i \sum_j p(b_j) \cdot p(a_i / b_j) \cdot \text{Log}_2 p(a_i / b_j), \quad (4)$$

де $p(a_i / b_j)$ – умовні ймовірності появи деякого стану a_i системи A за умови, що перед тим з'явився деякий стан b_j системи B . Аналогічно можна визначити умовну ентропію системи B щодо A , якщо в (4) поміняти місцями елементи a_i і b_j :

$$H(B/A) = \sum_i \sum_j p(a_i) \cdot p(b_j / a_i) \cdot \text{Log}_2 p(b_j / a_i) \quad (5)$$

У цьому випадку кількість інформації про системи A і B визначається на основі однієї із формул:

$$I(A, B) = k \cdot [H(A) - H(A/B)] = k \cdot [H(B) - H(B/A)], \quad (6)$$

де $H(A)$ і $H(B)$ – безумовні ентропії відповідно системи A і B , що визначаються як середня кількість інформації системи, яка припадає на один її стан. Якщо безумовні ентропії визначають середню корисну інформацію про систему, то умовні ентропії визначають втрати інформації за рахунок взаємного впливу однієї системи на другу (вплив завад). У випадку довільної кількості систем, що взаємодіють, для визначення кількості інформації, яку можна отримати про взаємозалежні системи, що взаємодіють, вводиться поняття ентропії вищого порядку, що зумовлено сумісною появою станів двох, трьох і більше систем перед появою стану розглядуваної системи. Отже, взаємна ентропія Z систем, що взаємодіють

$$H(A_1, A_2, \dots, A_Z) = \sum_{i_1} \sum_{i_2} \dots \sum_{i_Z} p(a_1, a_2, \dots, a_Z) \cdot \text{Log}_2 p(a_1, a_2, \dots, a_Z), \quad (7)$$

а умовна ентропія $Z-1$ порядку

$$H(A_Z / A_1, \dots, A_{Z-1}) = \sum_{i_1} \sum_{i_2} \dots \sum_{i_Z} p(a_1) \cdot \dots \cdot p(a_{Z-1}) \cdot p(a_Z / a_1, \dots, a_{Z-1}) \cdot \text{Log}_2 p(a_Z / a_1, \dots, a_{Z-1}).$$

Використовуючи співвідношення для умовних ймовірностей :

$$p(a_1, a_2, \dots, a_k / a_{k+1}, \dots, a_Z) = \frac{p(a_1, \dots, a_Z)}{p(a_{k+1}, \dots, a_Z)},$$

та враховуючи властивість ієрархічної мультиплікативності ймовірностей:

$$p(a_1, \dots, a_Z) = p(a_1) \cdot p(a_2 / a_1) \cdot p(a_3 / a_1, a_2) \cdot \dots \cdot p(a_Z / a_1, a_2, \dots, a_{Z-1})$$

маємо зв'язок між взаємною ентропією і безумовною та умовними ентропіями до Z-1 включно:

$$H(A_1, A_2, \dots, A_Z) = H(A_1) + H(A_2 / A_1) + H(A_3 / A_1, A_2) + \dots + H(A_Z / A_1, A_2, \dots, A_{Z-1}).$$

Записавши останню рівність ще раз з виділенням безумовної ентропії системи A_2 після перенесення всіх умовних ентропій з лівої сторони на праву і навпаки, отримуємо вираз для розрахунку взаємної кількості інформації довільної кількості взаємозалежних систем, що взаємодіють:

$$I(A_1, \dots, A_Z) = k \cdot [H(A_2) - H(A_2 / A_1) - H(A_3 / A_1, A_2) - H(A_4 / A_1, A_2, A_3) - \dots - H(A_Z / A_1, \dots, A_{Z-1})]$$

Зауважимо, що наведені співвідношення справедливі для дискретних просторів систем, що взаємодіють, використовуючи неперервні простори, слід оперувати не ймовірностями появи станів, а функціями густини ймовірностей станів і замість підсумування (формула (4), (5), (7)) здійснювати інтегрування.

Існує широкий клас задач, пов'язаних з побудовою інформаційних систем, до яких не можна безпосередньо застосувати наведені співвідношення для дискретних просторів.

До цих задач належать ті, в яких події на вході та виході є неперервними функціями часу в діапазоні їх передавання.

Якщо дискретні відліки таких сигналів беруть відповідно до теореми Котельникова, то втрат інформації в каналах зв'язку не буде, якщо не враховувати дії завад.

Розглядаючи взаємну інформацію неперервного простору станів елементів, до якої прямує взаємна інформація між скінченними областями, можна скористатися виразами для дискретних просторів, замінивши в них ймовірності на відповідні густини розподілу ймовірностей.

Густини ймовірностей неперервних просторів за заданої густини сумісного ансамблю АВ задаються рівностями

$$P_A(a_i) = \int_{-\infty}^{\infty} p(a_i, b_j) db_j,$$

$$P_B(b_j) = \int_{-\infty}^{\infty} p(a_i, b_j) da_i.$$

Середнє значення взаємної умовної інформації:

$$I(A, B / C) = \iiint p(a, b, c) \frac{\log_2 p(a, b / c)}{p(a / c) \cdot p(b / c)} da db dc.$$

Ентропія неперервних просторів може розглядатися як середнє значення власної інформації, аналогічно дискретним просторам:

$$H(A) = - \int_{-\infty}^{\infty} P_A(a) \cdot \log_2 P_A(a) da,$$

аналогічно умовна ентропія двох неперервних просторів, що взаємодіють

$$H(A / B) = - \iint P_{A,B}(a, b) \cdot \log_2 p(b / a) da db$$

і нарешті

$$I(A, B) = H(A) - H(A / B) = H(A) + H(B) - H(A, B).$$

Отже, співвідношення (6) є достатньо універсальним способом обчислення кількості інформації, незалежно від того, з яким простором станів маємо справу: неперервним чи дискретним. Відмінність лише в тому, що у випадку дискретного простору маємо справу з ймовірностями появи станів системи, а у випадку неперервних просторів – з функціями розподілу густини станів у просторі.

Алгоритмічний підхід

Принципово інший підхід до визначення кількості інформації запропонував російський вчений А. М. Колмогоров [4]. Його алгоритмічна теорія інформації ґрунтується на понятті складності алгоритму перетворення одного об'єкта на інший. За такого підходу принциповим є встановлення взаємного зв'язку між об'єктами, які досліджуються, та довжиною програми, що їх опрацьовує. Кількість інформації за теорією алгоритмів перетворення одних об'єктів на інші визначається як довжина програми, що забезпечує можливість перетворення об'єкта A на об'єкт B :

$$I = f[G(A, B)], \quad (8)$$

де G – програма перетворення об'єкта A на об'єкт B ; f – функція, що визначає довжину програми перетворення в бітах.

Зазначимо, що кількість інформації за такого підходу істотно залежить від вибору структурного елемента перетворення. Кількість інформації максимальна, якщо елементом перетворення вибрано піксель, а мінімальна у разі вибору в ролі елемента цілої літери, а також буде проміжною за вибору частини літери елементом перетворення. Виміряти ж кількість інформації для окремо взятої літери цілком неможливо, оскільки для реалізації алгоритму (8) необхідні два об'єкти. Порівнювати ж літеру з нею самою не має змісту, оскільки довжина програми перетворення у цьому випадку дорівнюватиме нулю.

Отже, за алгоритмічного підходу очевидні такі недоліки:

а) під час вимірювання кількості інформації як параметр використовується довжина програми перетворення, яка істотно залежить від структури елементів, що дають можливість перетворити один об'єкт на інший. Що дрібнішою є вибрана структура елементів, то довшою стає програма перетворення для тих самих об'єктів;

б) та сама програма перетворення може бути використана для опрацювання цілого набору об'єктів, аналіз яких потребує виконання різної кількості комп'ютерних команд, тоді як довжина програми залишається незмінною;

в) не визначений метод для вимірювання кількості інформації у разі розгляду окремого об'єкта, оскільки в алгоритмі вимірювання (8) їх потрібно два;

г) втрачається властивість адитивності інформації під час розгляду систем, що взаємодіють, а це істотно ускладнює їх аналіз.

Ці недоліки характерні й для семантичного підходу до вимірювання інформації, який ґрунтується на опрацюванні логічних тверджень:

$$I = \log_2 F(S), \quad (9)$$

де F – функція, що залежить від кількості станів у логічних твердженнях (змінних, фактах, правилах); S – логічне твердження або предикат.

Вимірювання кількості інформації у окремих літерах чи словах за підходу (9) неможливе в принципі, оскільки ні окремі літери, ні окремі слова не є логічними твердженнями.

Зазначених недоліків позбавлений підхід, який ґрунтується на застосуванні методів теорії розпізнавання і оперує із поняттям інформації та кількості інформації, які відмінні від класичних означень, отриманих на основі теорії ймовірності.

Образний підхід

Класична теорія інформації, яку розробив К. Шенон, сьогодні є досконалим універсальним апаратом розв'язання задач, пов'язаних з кодуванням інформації, її перетворенням та оптимальним передаванням по каналах зв'язку на великі відстані. Обмеження цієї теорії в тому, що вона ніяк не враховує семантику (спосіб утворення) інформації та повністю ігнорує людський фактор у формуванні інформації, тобто цілком ігнорується поняття цінності інформації.

Алгоритмічна теорія, що ґрунтується на понятті складності алгоритму, зробила крок у напрямі врахування способу утворення інформації. Незважаючи на недоліки цієї теорії, вказані у попередньому розділі, її основні положення використано в образній концепції теорії інформації [5, 6]. Згідно з цією концепцією під інформацією слід розуміти розпізнані образи, які зберігаються у пам'яті комп'ютера або будь-якої іншої кібернетичної машини. Образом вважається сигнал, записаний у сенсорну пам'ять сканувальних пристроїв кібернетичної машини. Отже, для отримання інформації слід реалізувати процедуру розпізнавання вхідного об'єкта – образу на основі його

відображення – певного еталона цього образу, який створений на основі домовленостей між відправником та приймачем повідомлень.

Кількість інформації I_o , що міститься в деякому образі, який отримала й успішно розпізнала кібернетична машина, визначається за формулою:

$$I_o = q^{-1}F[G(O)], \quad (10)$$

де q – ймовірність правильного розпізнавання образу; F – функція отримання довжини розгорнутої (з урахуванням циклів) програми розпізнавання образу; G – довжина програми розпізнавання образу, бітів.

Очевидно, за підходу (10) уможлиблюється обчислення кількості інформації складно-структурованих образів, які можуть бути окремими словами, реченнями або текстами чи малюнками. Кількість інформації істотно залежатиме від довжини розгорнутої програми розпізнавання, ймовірності правильного розпізнавання та функційних можливостей кібернетичної машини, яка реалізує процедуру розпізнавання. Для об'єктивного вимірювання кількості інформації згідно з формулою (10) слід враховувати, що:

а) програми розпізнавання образів мають бути оптимальними щодо власного розміру, швидкодії та функціональних можливостей;

б) збільшення кількості операцій (команд), які потрібно виконати машині для успішного розпізнавання, приводить до збільшення інформації, яку отримуємо від образу;

в) чим меншою є ймовірність правильного розпізнавання образу, тим більшу кількість інформації він містить;

г) кількість інформації у незнаковому повідомленні дорівнює довжині цього повідомлення, вираженого в бітах, помноженій на довжину програми розпізнавання одного біта цього повідомлення і розділеній на ймовірність правильного розпізнавання повідомлення.

Перші експерименти [6] підтвердили право на існування такого підходу до визначення кількості інформації, хоча його ефективність може виявитися лише в процесі постановки численних експериментів з розпізнаванням різноструктурованих об'єктів. Недоліки цього підходу такі:

а) значною мірою суб'єктивна оцінка кількості інформації, яку містить образ, що зумовлено як технічними можливостями кібернетичної машини, так і якістю програми розпізнавання;

б) втрачена властивість адитивності кількості інформації, наявна у класичній теорії інформації;

в) надмірна громіздкість формули (10), що ускладнює проведення аналітичних оцінок.

Для усунення другого недоліку видається доцільним розрахунок кількості семантичної (образної) інформації виконати у логарифмічному масштабі:

$$I_o = -\log_2 q + \log_2 f(O) \quad (11)$$

де f – довжина програми розпізнавання образів, виражена у кількості операцій (машинних команд), необхідних для успішного розпізнавання образу.

Підхід до вираховування цінності інформації

Стосовно цінності інформації, то про неї можна говорити за потреби досягнення певної мети після того, як користувач отримує інформацію, тобто забезпечення досягнення деякої цільової функції. У роботі [5] А. А. Харкевич запропонував цінність інформації обчислювати як

$$F = \log_2 \frac{p_0}{p_1} \quad (12)$$

де p_0 – ймовірність правильного розв'язання проблеми до отримання інформації; p_1 – ймовірність правильного вирішення проблеми після отримання інформації. Такий підхід має право на існування, хоча і викликає сумніви його ефективність та доцільність застосування. По-перше, одиницею вимірювання цінності інформації за такого підходу є біт, як і у випадку обчислення кількості інформації. Очевидно, введення нової величини потребує нової розмірності або розрахунки треба виконувати у безрозмірних одиницях. По-друге, формула (12) не може претендувати на об'єктивність, оскільки оцінки p_0 і p_1 здійснюватиме користувач. По-третє, методика оцінювання цих ймовірностей не є очевидною. По-четверте, цінність інформації є динамічною величиною [7] і в міру надходження інформації мала би змінюватися. Очевидно,

обчислюючи цінність інформації, цільовою функцією для досягнення поставленої мети слід вибрати задоволення певних потреб користувача (матеріальних, духовних, естетичних, смакових, пізнавальних та інших) або виконання певних дій.

Найдоцільнішим з урахуванням розмаїття потреб користувача видається підхід, за якого цінність інформації розраховуватимемо у відсотках: 100 % – за умови цінності інформації; 0 % – за умови, що цільова функція не досягнута. Отже, цінність інформації F для сформованого поточного значення цільової функції Z та досягнутої ефективності E після отримання повідомлення на певний момент часу можна задати у вигляді правила:

$$F = \frac{100\%, \text{ якщо } Z - E = 0}{0\%, \text{ якщо } Z - E > 0} \quad (13)$$

Для підвищення точності вирахування цінності інформації доцільно мати сформовану цільову функцію і на наступні моменти часу та розширити шкалу розрахунку цінності інформації з певним кроком. Перша умова істотно підвищує цінність інформації після отримання всього повідомлення, а друга може бути реалізована за формулою:

$$F = \frac{E}{Z} \cdot 100 \% \quad (14)$$

Формула (13) є уточненням (12), що підтверджено розглядом конкретних прикладів із систем масового обслуговування та теорії ігор.

Очевидно, для реалізації описаного підходу щодо обчислення цінності інформації доцільно використовувати декларативні мови програмування (Лісп, Пролог або їх модифікації залежно від конкретної предметної задачі) [26–34], які найвдаліше пристосовані для реалізації логічних функцій виду (10)–(14) і дають змогу розв’язувати задачі, пов’язані з якісним розпізнаванням та аналізом об’єктів складної структури (розпізнавання почерку, рукописного тексту, психофізіологічного стану особи, побудова та аналіз систем зберігання, опрацювання та захисту інформації) [11–25, 36] та реалізації відповідних до конкретної прикладної чи наукової задачі цільових функцій.

Висновки і перспективи подальших наукових розвідок

У роботі автори висвітлили три класичні підходи до вирахування кількості інформації та оцінки її цінності: класичний, алгоритмічний та образний. Подано порівняльну характеристику наведених підходів, вказано переваги та обмеження кожного з них та визначено перспективи їх використання.

Автори запропонували новий підхід до оцінки цінності інформації на основі образного підходу, який розширює область його застосування та може бути успішно реалізований з використанням декларативних мов програмування [26–34] або універсальних мов моделювання (UML) [35, 37].

Подальшого успіху у використанні запропонованого підходу можна досягти, здійснивши статистичні дослідження конкретних прикладних задач, пов’язаних з необхідністю оцінки як кількості, так і цінності, отриманої унаслідок застосування інформації.

Запропонований підхід видається доцільним для прикладного застосування, оскільки сприятиме розвитку як методів розпізнавання, так і власне теорії інформації.

1. Бриллюэн Л. Наука и теория информации / Л. Бриллюэн. – М.: Госиздат, 1960. – 392 с.
2. Хартли Р. В. Теория информации и её приложения / Р. В. Хартли. – М.: Физматгиз, 1959. – 356 с.
3. Шеннон К. Работы по теории информации и кибернетике / К. Шеннон. – М.: Изд-во иностран. л-ры, 1963. – 286 с.
4. Колмогоров А. Н. Три подхода к определению понятия “количество информации” / А. Н. Колмогоров // Проблемы передачи информации. – Т.1. – Вып. 1. – 1965. – С. 63–67.
5. Харкевич А. А. О ценности информации / А. А. Харкевич // Проблемы кибернетики. – Вып. 4. Физматгиз. – 1960. – С. 53–57.
6. Партико З. В. Образна концепція теорії інформації / З. В. Партико. – Львів: Вид-во ЛНУ ім. І. Франка, 2001. – 98 с.
7. Бонгард М. М. О понятии “полезная информация” / М. М. Бонгард // Проблемы кибернетики. Вып. 8. – Физматгиз. – 1963. – С. 71–102.
8. Теслер Г. С. Новая кибернетика / Г. С. Теслер. – К.: Логос, 2004. – 404 с.
9. Сявавко М. С. Ентропія як показник розмитості нечіткої множини / М. С. Сявавко, М. І. Рожанківська // 36.

наукових праць ЛДІНТУ ім. В. Чорновола. – № 2. – 2009. – С. 3–19. 10. De Luca A. On the convergence of entropy measures of fussy sets / A. De Luca, S. Termini // *Kybernetes*. – Vol. 6. – 1971. – P. 219–227. 11. Шустер Г. Детерминированный хаос: введение. / Г. Шустер; пер. с англ. – М.: Мир, 1988. – С. 240. 12. Заяць В. М. Математичний опис системи розпізнавання користувача комп'ютера / В. М. Заяць, М. М. Заяць // Зб. “Фізико-математичне моделювання та інформаційні технології”. – Львів. – 2005. – Вип. 1. – С. 146–152. 13. Харкевич А. А. Опознавание образов / А. А. Харкевич // *Радиотехника*. – 1959. – Т. 14. – С. 15–19. 14. Фукунага К. Введение в статистическую теорию распознавания / К. Фукунага. – М.: Наука, 1979. – 512 с. 15. Горелик А. Л. Методы распознавания / А. Л. Горелик, В. А. Скрипник. – М.: Высшая школа, 1989. – С. 232. 16. Дуда Р. Распознавание образов и анализ сцен. / Р. Дуда, П. Харт. – М.: Мир, 1976. – С. 512. 17. Заяць В. М. Визначення пріоритету детермінованих ознак при побудові системи розпізнавання об'єктів / В. М. Заяць, О. Шокира // Зб. праць науково-практичної конф. ЛДІНТУ імені В. Чорновола “Математичне моделювання складних систем”. – 2007. – С. 135–137. 18. Заяць В. М. Архітектура подіє-орієнтованих систем на прикладі системи розпізнавання рукописного тексту / В. М. Заяць, Д. О. Іванов // Вісник Нац. ун-ту “Львівська політехніка” “Комп'ютерна інженерія та інформаційні технології”. – Львів. – 2004. – № 530. – С. 78–83. 19. Томашевський О. М. Методи та алгоритми системи захисту інформації на основі нейромережових технологій: автореф. дис. ...канд. техн. наук: 05.13.23 / О. М. Томашевський – Львів, 2002. – С. 20. 20. Чалая Л. Є. Сравнительный анализ методов аутентификации. *Computer journal of Man-Machine Studies*. – 1988. – Vol. 28, n. 1. – P. 67–76. 21. Cohell O. Biometric Identification System Based in Keyboard Filtering / O. Cohell, J. Badia, G. Torres // *Proc. Of XXXIII Annual IEEE International Carnahan Conference of Security Technology*. – 1999. – P. 203–209. 22. Вовк О. Б. Проблеми захисту шрифтів як специфічних об'єктів авторського права / О. Б. Вовк // Вісник Нац. ун-ту “Львівська політехніка” “Інформаційні системи та мережі”. – 2008. – № 610. – С. 85–83. 23. Платонов А. В. Використання експертних ситуативних моделей у сфері державної безпеки / А. В. Платонов, І. В. Баклан, К. В. Крамер // Зб. праць міжнар. наукової конф. ISDMCI' 2008. – Т. 1. – Єваторія, 2008. – С. 39–43. 24. Суздальцев А. И. Определение психофизического состояния оперативного персонала по клавиатурному почерку на нефтеперерабатывающих мини-заводах / А. И. Суздальцев, В. А. Лобанова, В. Г. Абаишин // *Нефтегазовое дело*. – 2006. – С. 1–6. 25. Zayats V. Structural method of hand-written text recognition / V. Zayats, D. Ivanov // *Pros. International Conf. “The experience of designing and application of CAD systems in microelectronics”*. – Lviv-Polyana, 2005. – P. 493–494. 26. McCarthy J. Recursive functions of symbolic expressions and their computation by machine // *Comm. ACM*. – 1960. – Vol. 3. – P. 184–195. 27. Бадаєв Ю. І. Теорія функціонального програмування. Мови CommonLisp та AutoLisp. – К., 1999. – 150 с. 28. Заяць В. М. Функційне програмування: навч. підручник. – Львів: Бескид Біт, 2003. – 160 с. 29. Хендерсон П. Функциональное программирование. Применение и реализация. – М.: Мир, 1983. – 349 с. 30. Братко И. Программирование на языке Пролог для искусственного интеллекта. – М.: Вильямс, 2004. – 640 с. 31. Ин Ц., Соломон Д. Использование Турбо-Пролога; пер. с англ. – М.: Мир, 1993. – 608 с. 32. Заяць В. М. Логічне програмування: Ч. 1: Конспект лекцій з дисципліни “Логічне програмування” для студентів базового напрямку 6.08.04 “Програмне забезпечення автоматизованих систем”. – Львів: Вид-во Нац. ун-ту “Львівська політехніка”, 2002. – 48 с. 33. Макаллистер Дж. Искусственный интеллект и Пролог на микро ЭВМ. – М.: Машиностроение, 1990. – 240 с. 34. Заяць В. М., Заяць М. М. Логічне і функційне програмування: навч. посіб. – Львів: Бескид-Біт, 2006. – 352 с. 35. Erich Gamma, Richard Helm, Ralih Johnson, John Vlissides. *Wzorce projektowe*, 2012. 36. Nilsson N. J. *Principles of Artificial Intelligence*. – Tioga– Springer–Verlag, 1980. – 164 с. 37. Stanislaw Wrycha i inni. *UML 2.1. Cwiczenia*.