

¹Харківський національний університет радіоелектроніки, Харків, Україна²Харківський навчально-науковий інститут ДВНЗ Університет банківської справи, Харків, Україна

СТАТИСТИЧНІ РОЗПОДІЛИ ТА ЛАНЦЮЖКОВЕ ПОДАННЯ ДАНИХ ПРИ ВИЗНАЧЕННІ РЕЛЕВАНТНОСТІ СТРУКТУРНИХ ОПИСІВ ВІЗУАЛЬНИХ ОБ'ЄКТІВ

Предметом досліджень статті є моделі для встановлення рівня релевантності зображень у просторі розподілів для дескрипторів ключових точок при розпізнаванні візуальних об'єктів у системах комп'ютерного зору. **Метою** є створення методу структурного розпізнавання зображень на підставі впровадження ланцюжкових моделей даних із використанням ймовірнісних розподілів множини дескрипторів. **Завдання:** розроблення математичних та програмних моделей для ефективного за швидкодією аналізу даних при визначенні релевантності структурних описів, вивчення властивостей, атрибутів застосування, значень параметрів цих моделей, оцінювання результативності за наслідками оброблення конкретних зображень. Застосовуваними **методами** є: детектор BRISK для формування дескрипторів ключових точок, апарат інтелектуального аналізу даних, методи побітового оброблення та побудови розподілів бітових даних, апарат метричного визначення релевантності, програмне моделювання. Отримані такі **результати**. Перехід від опису множин дескрипторів до ймовірнісних розподілів фрагментів і зіставлення образів у просторі розподілів забезпечують необхідну результативність розпізнавання. Оброблення та аналіз даних виконується у сотні разів швидше, ніж традиційний підрахунок голосів. Оброблення та аналіз сполучень бітів формує значимі властивості для сукупності елементів опису зі збереженням структури даних і їх уніфікації. Зі збільшенням числа бітів у фрагменті розподілу зростає відстань між зображеннями, що сприяє збільшенню ступеня їх розрізнення. Ланцюговим поданням та застосуванням розподілів створюється новий простір даних, що дає можливість суттєво покращити показники функціонування систем розпізнавання зображень. **Висновки.** Наукова новизна дослідження полягає в удосконаленні методу структурного розпізнавання зображень на основі впровадження узагальненої ланцюгової структури опису із використанням значень розподілу для фрагментів множини дескрипторів ключових точок, що змістовно відображають властивості зображень об'єктів і забезпечують результативне розпізнавання. Практична значущість – досягнення суттєвого рівня підвищення швидкодії обчислення релевантності, підтвердження результативності запропонованих модифікацій на прикладах зображень, отримання прикладних програмних моделей для дослідження та впровадження методів класифікації в системах комп'ютерного зору.

Ключові слова: структурні методи розпізнавання зображень, ключова точка, детектор BRISK, ланцюгове подання, розподіл даних фрагменту, спільний дескриптор, релевантність описів, голосування, манхеттенська метрика, швидкодія визначення релевантності.

Вступ

Статистичні моделі набули істотного поширення в інтелектуальних методах аналізу даних, в тому числі і при прийнятті класифікаційних рішень у системах комп'ютерного зору [1-4]. Їх ключовою перевагою є використання для класифікації певної узагальненої інформації про властивості класів розпізнаваних об'єктів, що дає змогу результативніше врахувати особливості об'єктів у просторі ознак. Найбільш точно характеристики об'єкту відображаються при безпосередньому використанні розподілів даних [2,3]. Як правило, дані у комп'ютерному зорі мають багатовимірний ймовірнісний розподіл. Але іноді їх можна подати у вигляді множини розподілів

одновимірних величин, що значно спрощує класифікацію та знижує обчислювальні затрати [4].

Методи розпізнавання зображень візуальних об'єктів за множиною ключових точок (КТ) отримали практичне застосування із-за таких важливих властивостей, як інваріантність до геометричних перетворень об'єктів, можливість розпізнавання в умовах часткового подання, стійкість до різноманіття завад [4-6].

Аналіз сигналів шляхом представлення їх як множини фрагментів досить популярний у методах комп'ютерного зору [7-10], так як часто найважливіша інформація про образи об'єктів зосереджується у окремих деталях.

Бітова природа дескрипторів КТ у просторі B^n бінарних векторів (n – ступінь двійки) дає

можливість впровадити подання та аналіз дескриптора як ланцюжка елементів (наприклад, байтів), діапазон значень яких відомий. Це дає змогу здійснювати статистичний аналіз даних з урахуванням внутрішнього змісту наявного дескрипторного опису об'єкту. З точки зору важливості інформації всі елементи рівноцінні, але місце їх розміщення у складі дескриптора фіксоване, тому є можливість аналізувати чи обробляти упорядковані послідовності елементів. Ланцюжкова структура допускає застосування підходів інтелектуального аналізу, заснованих на ймовірнісних оцінках наявних значень даних, щоб прийняти рішення про віднесення об'єкту з описом у вигляді множини дескрипторів до відповідного класу.

Статистичне подання є одним із найбільш популярних інструментів у сучасному інтелектуальному аналізі даних задля виявлення закономірностей чи встановлення системи знань, що містять дані [3, 7-11]. Воно ґрунтується на виявленні у змісті аналізованої інформації таких суттєво значимих характеристик, вартість яких оцінюють параметром частоти зустрічальності як оцінки ймовірності. Навчання та самонавчання в основному теж базуються на визначенні найчастіше вживаних компонентів аналізованих даних [4-6]. На підставі аналізу скінченного об'єму даних здійснюють апроксимацію ймовірнісного розподілу, що дає змогу розпізнавати їх образи. Важливим аспектом, що досить точно відображає властивості чи характеристики даних, є безпосереднє використання значень розподілів у методах розпізнавання [2].

Один із підходів заснований на використанні таких статистичних характеристик, як математичне очікування, дисперсія, оцінки медіани та ін. Але більш інформаційним є все-таки використання безпосередньо самих значень розподілів, як це робиться в методах інтелектуального аналізу. Це дає змогу більш чуттєво врахувати відмінності та особливості значень даних, що відображають властивості розпізнаваних класів.

Метою статті є опрацювання методу структурного розпізнавання зображень на підставі впровадження ланцюжкових моделей даних із використанням ймовірнісних розподілів множини дескрипторів, що формують структурний опис.

Задачами дослідження є розроблення математичних та програмних моделей ефективного за швидкодією оброблення даних при обчисленні релевантності структурних описів, вивчення властивостей та атрибутів застосування цих моделей з урахуванням значень їх параметрів, оцінювання ефективності за результатами оброблення конкретних зображень.

Формування ланцюгової структури дескриптора КТ

Зважаючи на те, що параметр n розміру аналізованих елементів даних сягає кількох сотень (для дескриптора BRISK $n = 512$), побудова, аналіз, та оброблення такого роду багатовимірних розподілів викликає труднощі для обчислень, доцільно впровадити узагальнюючий статистичний аналіз для сукупностей фрагментів даних, що подаються у вигляді ланцюжка. Одним із варіантів оброблення є розбиття кожного дескриптора КТ розміру n на $m \ll n$ непересічних фрагментів, що повністю складають початковий розмір n .

Розглянемо подання структури дескриптора КТ зображення у вигляді послідовності (ланцюжка) із m елементів фіксованого розміру (рис. 1).

елемент 1	елемент 2	...	елемент m
-----------	-----------	-----	-------------

Рис. 1. Ланцюжкова структура дескриптора даних

Кожний елемент приймає значення із фіксованого діапазону. Наприклад, якщо елемент – 1 біт, то маємо два значення (0, 1), якщо байт – множину значень $0, \dots, 255$. Опишемо параметром l кількість значень (ланок розподілу), що приймає елемент. Дескриптор BRISK [5] із 512 бітів тепер може бути представлений різноманітним варіантом ланцюжка із $m = 1, 2, \dots, 512$ елементів. У байтовому поданні значення $m = 512/8 = 64$.

Для розпізнаваного візуального об'єкту визначимо його дескрипторний опис у вигляді скінченної множини $Z = \{z_v\}_{v=1}^s$, $z_v \in B^n$ із s бінарних дескрипторів КТ (наприклад, ORB, BRISK). У результаті ланцюжкового подання опис Z буде мати вид матриці із s рядків по m елементів у рядку.

Статистичне подання множини дескрипторів об'єкта

Постає задача побудувати таке відображення $Z \rightarrow Q$ із множини бінарних векторів – дескрипторів КТ у множину значень Q розподілів значно меншої потужності w , що дають можливість проводити ідентифікацію та розрізнення візуальних об'єктів за їх описом Z . Розподіл $q \in Q$ подамо у вигляді вектора $q = \{q_1, \dots, q_w\}$ цілих чисел, сума яких дорівнює загальному об'єму s аналізованих даних опису, тобто $\sum_i q_i = s$.

Для кожного з m фрагментів на підставі множини Z побудуємо розподіл $q = \{q_1, \dots, q_w\}$, причому величина w цілком визначається

діапазоном значень даних для фрагмента. Наприклад, для дескриптора BRISK при розбитті на байти для $n=512$ маємо $m=64$, $w=256$. Повністю множина Z буде описана системою розподілів $Q=\{q^v=(q_1, \dots, q_w)^v\}_{v=1}^m$ для кожного із m фрагментів.

Розглянемо приклад аналізу множини Z за змістом окремих байтів: для кожного з m байтів з фіксованим розміщенням у складі дескриптора підрахуємо кількість q_i входження його фіксованого значення i до загального змісту із s даних, $q_i \leq s$. Значення $p=q_i/s$ є частотою (оцінкою ймовірності) його появи у множині Z .

Як результат аналізу за байтами маємо табл. 1 цілих чисел q_i , що містить частотні оцінки для різноманіття можливих значень байта, причому загальна сума величин q_i дорівнює s : $\sum_i q_i = s$.

Таблиця 1. Розподіл значень байту за описом Z

0	1	2	...	254	255
q_0	q_1	q_2	...	q_{254}	q_{255}

У другому рядку табл. 1 маємо характеристики розподілу всіх можливих значень байта на підставі розгляду конкретної множини КТ, що складає опис Z . Рядок відображає характеристики ймовірнісної структури опису у розрізі байтового подання.

Загалом на підставі множини Z отримуємо матрицю даних $Q=\{q_{vw}\}$ із 256 рядків, $k=0, \dots, 255$, та m стовпців. Матриця Q містить оцінки ймовірностей щодо появи конкретних значень у фрагментах розміром в байт для множини Z . Зауважимо, що при аналізі ланцюжків у формі бітів достатньо використовувати тільки одну із ланок розподілу замість двох. Наявність умови $\sum_i q_i = s$ дає можливість контролювати вагомість рівня відповідності двох розподілів, так як діапазон їх значень заданий. На цьому факті можна ґрунтувати узагальнену оцінку рівня значимості щодо еквівалентності розподілів.

Встановлення релевантності описів на підставі значень розподілів

Загальний підхід зводиться до зіставлення усіх наявних значень матриць $Q=\{q_{vw}\}$, отриманих у межах опису окремих об'єктів. Матриця Q відображає статистичні властивості опису як розподіл значень його сегментів.

Релевантність обчислюємо у вигляді відстані між матрицями для різних описів. Визначимо релевантність r описів a та b на підставі зіставлення множини їх розподілів щодо введеної

системи фрагментів як манхеттенську відстань між матрицями $Q(a)$, $Q(b)$:

$$r[Q(a), Q(b)] = \sum_{v=1}^m \sum_{w=1}^l |q_{vw}(a) - q_{vw}(b)| \quad (1)$$

Зауважимо, що тут є можливість дію додавання за зовнішньою сумою в (1) зупиняти у процесі обчислень, якщо контрольоване значення $r \geq \delta_r$ перевищить деякий поріг δ_r , фіксуючи певну відсутність подібності описів [7]. Важливо, що у виразі (1) обчислення виконуються виключно із використанням цілих чисел, а значення r у ході обчислень тільки зростає із-за накопичення додатних значень.

Відмітимо, що умова використання тільки цілих значень розподілів ґрунтується на однаковому значенні s для потужностей описів. Ця умова легко досягається на попередньому етапі формування опису. Якщо ж із-за якихось вимог прикладних задач це обмежено, то треба перейти до оброблення нормованих даних.

При обчисленні міри (1) можуть бути також застосовані інші схеми логічного аналізу, пов'язані з прийняттям рішення по заданому числу δ_m суттєво подібних розподілів фрагментів. Наприклад, якщо число схожих фрагментів уже перевищує значення δ_m , то описи визнаються релевантними [7]. При цьому аналогічно моделям зіставлення дескрипторів КТ [4, 6] необхідно також ввести поріг δ_q для визнання еквівалентності двох розподілів. Зважаючи на те, що у мірах типу (1) застосовуються виключно цілі числа, такі пороги також приймають цілі значення та можуть бути адаптовані до бази розпізнаваних зображень.

Метрика (1) реалізує порівняння розподілів, а не їх параметрів чи характеристик. Це найбільш точніше зіставлення із використанням повної інформації, що націлене на виявлення навіть незначних відмінностей у описах. Кількість l ланок розподілів для розміру дескриптора 512 в залежності від числа бітів у фрагменті та від кількості m фрагментів даних наведена в табл. 2.

Таблиця 2. Число ланок розподілу в залежності від параметрів фрагмента

Число бітів фрагмента	1	2	4	8	512
m	512	256	128	64	1
l	2	4	16	256	2^{512}

Зважаючи на те, що при визначенні релевантності описів у побудованому просторі значень розподілів даних загальний обсяг обчислень прямо пропорційний добутку $m \cdot l$ числа фрагментів та числа ланок, із аналізу значень табл. 2

маємо висновок, що для вирішення цієї задачі з найменшими витратами обчислень треба прагнути до застосування параметрів лівої половини табл. 2. З іншого боку, інформації від розподілів окремих бітів чи їх пар може не вистачити для забезпечення необхідного рівня розрізненості описів, і тоді прийдеться йти на збільшення обсягу обчислень.

Одним із шляхів спрощення аналізу та оброблення ланцюжкових даних є використання значень окремих параметрів побудованих розподілів. Процедури визначення спільного дескриптора (СД) [4, 6] для побудови концентрованого опису зображення на підставі введеного ланцюжкового подання можна базувати, наприклад, на застосуванні найбільш вживаних шаблонів – бінарних масок, узорів, паттернів (а фактично – цілих чисел), що найчастіше зустрічаються у межах опису.

Визначимо максимальне значення у кожному стовпці матриці Q і сформуємо ланцюжок СД d як результат зчеплення найбільш вживаних у межах опису байтів

$$d_j = \arg \max_{i=1, \dots, s} q_i, j = \overline{1, m}, \quad (2)$$

$$d = d_1 \& d_2 \& \dots \& d_m,$$

а значення d_j як аргумент величини q_i отримано для кожного з байтів рядка табл. 1. Таким чином, СД для опису Z тут має вигляд вектору із m байтів, що найчастіше зустрічаються у відповідному місці кожного із елементів множини дескрипторів.

У цій спрощеній моделі важливо враховувати лише достатньо вагомий за максимумом величини q_i характеристикою значимості $q^* = \max_{i=1, \dots, s} q_i$, наприклад, за умови $q^* \geq 0,5$ або $q^* \geq 0,25$. Можна на попередньому етапі проаналізувати і сформувати список байтів, які для конкретного еталону мають досить високий рівень підтримки q^* всередині опису. Решту байтів виключаємо з аналізу. У випадку, коли аналіз здійснюється за окремими бітами (замість байтів), то перевірка значимості процедурою (2) реалізується автоматично, так як у табл. 1 буде лише два стовпці. Обмеження на значення підтримки забезпечує адаптацію створеного СД до аналізованого масиву даних і точніше відтворення його змісту. Визначення релевантності СД для цього способу виконується шляхом обчислення відстані між векторами або з врахуванням по-елементного аналізу. Аналіз виду (2) відповідає принципу максимуму апостеріорної ймовірності. Для суттєво різних описів такий аналіз приносить необхідний результат [6].

Подання та аналіз даних у вигляді множини фрагментів отримали значне поширення у теорії

комп'ютерного зору [7, 8]. Воно забезпечує можливість прийняття рішень за окремими елементами візуальних об'єктів, що важливо у прикладному сенсі. При цьому керування такими характеристиками фрагментної системи, як значення порогів на кількість та подібність фрагментів, покладається на користувача. Зрозуміло, що встановлення оптимального розміру елемента, порогу для визначення еквівалентних за описом елементів, порогу для встановлення значущості опису елемента повинно бути виконано, виходячи із заданої бази зображень, в межах якої виконується розпізнавання.

Ланцюжкове подання розкриває шлях для побудови множини моделей класифікації на підставі системи фрагментів [10]. Кожен з m різних «маленьких» еталонів незалежно голосує за відповідний клас. Так вводиться структура всередині системи дескрипторів, де окремі частини підтримують або не підтримують пропонуване рішення.

Оброблення множини дескрипторів у пропонованій схемі оброблення може бути виконане в інших варіантах аналізу груп байтів, наприклад, пар, трійок, а також окремих бітів. Не виключений також і спосіб створення покриття дескриптора, тобто з урахуванням можливого перетину змісту аналізованих елементарних структур даних.

Результати комп'ютерних експериментів

Нами проведено програмне моделювання розроблених методів мовою C# у середовищі Visual Studio 2017 із використанням засобів бібліотеки Open CV [12]. Формування дескрипторів BRISK, побудова розподілів ланцюжкового подання з параметрами 1-го та 2-х бітів, а також визначення СД здійснювалося на зображеннях ікон розміром 400x540 (рис. 2).



Рис. 2. Зображення ікони «Казанська Богоматір»

Число 700 дескрипторів КТ вибрано однаковим для зображень. Релевантність описів обчислювалась як відстань (1) між розподілами, як відстань Хемінга між побудованими бінарними СД та як нормоване число голосів еквівалентних значень дескрипторів КТ. Міра подібності за голосуванням (підрахунок подібних дескрипторів з порогом еквівалентності 20% від максимуму відстані) для двох різних ікон склала 0,12 (максимум – 1), в той час як нормована до числа КТ відстань Хемінга для їх СД, побудованих за системою фрагментів, дорівнює 0,068 (мінімум 0).

Приклади розподілів значень описів у вигляді множин дескрипторів КТ для окремих бітів та пар бітів наведені у табл. 3, 4.

Таблиця 3. Число нулів у бітовому розподілі для перших 8-ми бітів

Біт№	1	2	3	4	5	6	7	8
Ікона 1	489	490	314	449	209	183	347	136
Ікона 2	506	484	277	438	206	194	358	140

Таблиця 4. Приклади розподілів для пар бітів

П а ра №	Ланки для ікони 1				Ланки для ікони 2			
	00	01	10	11	00	01	10	11
1	468	21	22	189	459	47	25	169
2	213	101	236	150	175	102	263	160
3	139	70	44	447	144	62	50	444
4	77	270	59	294	80	278	60	282
5	86	36	106	472	87	41	50	522
6	72	213	50	365	74	203	61	362
7	91	39	93	477	92	28	45	535
8	139	113	149	299	96	116	161	327

Як можна побачити із цих таблиць, розподіли для двох ікон достатньо відрізняються, що підкреслює результативність запропонованого методу в аспекті якості розрізнення. Відстань (1) уже для 8-ми бітового фрагмента опису із табл. 3 складає величину 200, максимум відстані при цьому дорівнює $1400 \times 8 = 11200$. Сумарна відстань досягла 43653 (6,1% від досяжного максимуму). Відстань (1) для першої пари ланок розподілу табл. 4 дорівнює

58, для другої пари – 76, для третьої – 22. Загальна відстань при цьому складає величину 30822 (8,6% від досяжного максимуму). Зважаючи на те, що мінімальне значення (1) при цьому дорівнює нулю, можна вказати на наявні чіткі позитивні властивості розподілів щодо розрізнення описів. Це ж саме відноситься і до відстані між побудованими СД обох досліджуваних зображень.

Зауважимо також, що зі збільшенням кількості бітів у сформованому розподілі зростає у відсотковому відношенні і відстань між зображеннями ікон, що сприяє збільшенню ступеня їх розрізнення. Цей факт спостерігався при різному числі дескрипторів у аналізованих описах. У той же час зі зменшенням числа КТ у описах відстань у просторі розподілів теж зростає.

Аналіз обчислювальних витрат для розроблених моделей аналізу даних показав, що час обчислення для модифікацій на базі розподілів у порівнянні з традиційним голосуванням зменшився приблизно у 1000 разів. Конкретно значення склали 8 мсек для побітового аналізу і 11 мсек для аналізу пар бітів, в той час як для голосування знадобилося майже 11 сек. При цьому виграш у швидкодії зростає пропорційно числу використовуваних КТ. Моделі аналізу за одним та парою бітів дають приблизно однакові показники швидкодії, але зрозуміло, що зі зростанням числа ланок розподілу час обчислень дещо збільшується. Загалом, розроблені методи забезпечують значно вищу швидкодію функціонування системи розпізнавання.

Висновки

Перехід від множин дескрипторів КТ до ймовірнісних розподілів для фрагментів цих дескрипторів і зіставлення образів у просторі розподілів забезпечують необхідну результативність розпізнавання. Таке оброблення виконується значно швидше, ніж традиційний підрахунок голосів КТ.

Оброблення за аналізом значень бітів чи їх сполучень формує значимі властивості для сукупності елементів опису, а послідовне об'єднання їх у новий вектор зберігає вихідну структуру даних та їх уніфікацію. Найбільш детальний аналіз зводиться до побітового.

Зі збільшенням кількості бітів у фрагменті сформованого розподілу зростає у відсотковому відношенні і відстань між зображеннями, що сприяє покращенню якості їх розрізнення.

Ланцюговим поданням та застосуванням розподілів даних створюється новий простір для аналізу даних, що суттєво покращує показники функціонування систем розпізнавання зображень.

Практичні рекомендації із проведеного дослідження полягають у застосуванні однобітового аналізу як оптимального з точки зору

швидкодії обчислень. Якщо показників достовірності розпізнавання при цьому не досягнуто, виникає необхідність оброблення та аналізу пар чи груп бітів.

Наукова новизна дослідження полягає в удосконаленні методу структурного розпізнавання зображень на основі впровадження ланцюгової структури опису із використанням значень розподілу для фрагментів множини дескрипторів ключових точок, що змістовно відображають властивості зображень об'єктів.

Практична значущість роботи – досягнення суттєвого рівня підвищення швидкодії обчислення релевантності, підтвердження результативності запропонованих модифікацій на прикладах зображень, отримання прикладних програмних моделей для дослідження та впровадження методів класифікації в системах комп'ютерного зору.

Перспективи дослідження можуть бути пов'язані із визначенням оптимальних параметрів фрагментації даних задля забезпечення потрібного рівня результативності системи розпізнавання.

Список літератури

1. Szeliski R. Computer Vision: Algorithms and Applications / R. Szeliski. – London: Springer, 2010. – 979 p.
2. Gadetska S.V. Statistical Measures for Computation of the Image Relevance of Visual Objects in the Structural Image Classification. // *Володимир Г. ГОРОХОВАТСЬКИЙ* Gadetska, V.O. Gorokhovatskyi, K. G. Solodchenko // *Виктор Г. ГАДЕЦЬКА* Systems and Radio Engineering. – 2018. – №2 (48). – С. 1041–1053. *Володимир Г. ГОРОХОВАТСЬКИЙ*
3. Han, J. Data Mining: Concepts and Techniques / J. Han, M. Kamber. – Amsterdam e.a.: Morgan Kaufmann Publishers, 2006. – 754 p.
4. Gorokhovatskyi V.A. Image Classification Methods in the Space of Descriptors of a Set of the Key Point Descriptors / V.A. Gorokhovatskyi // *Володимир Г. ГОРОХОВАТСЬКИЙ* Telecommunications and Radio Engineering. – 2018, 77 (9), pp. 787-797.
5. Stefan Leutenegger, Margarita Chli, Roland Y. Siegwart. BRISK: Binary Robust Invariant Scalable Keypoints. – Computer Vision (ICCV), pp. 2548 – 2555, 2011.

Автори

Моб. тел. для зв'язку : 097-33-44-538, **Гороховатський Володимир Олексійович**

ГОРОХОВАТСЬКИЙ
Володимир Олексійович
(GOROKHOVATSKYI
Volodymyr)

Харківський національний університет радіоелектроніки, Харків, Україна,
доктор технічних наук, професор, професор кафедри інформатики,
e-mail: gorohovatsky.vl@gmail.com; ORCID: 0000-0002-7839-6223

ГАДЕЦЬКА
Світлана Вікторівна

Харківський навчально-науковий інститут ДВНЗ Університет банківської справи, Харків,
Україна,

6. Гороховатський В.О. Застосування апарату аналізу та оброблення бітових даних у методах класифікації зображень за множиною ключових точок / В.О. Гороховатський, К. Г. Солодченко // Системи управління, навігації та зв'язку. – 2018. – №2 (48). – С. 63–67.

7. Баклицкий В.К. Методы фильтрации сигналов в корреляционно-экстремальных системах навигации/ В.К. Баклицкий, А.М. Бочкарев, М.П. Мусьяков. – М.: Радио и связь, 1986. – 216 с.

8. Путятин Е.П. Обработка изображений в робототехнике / Е.П. Путятин, С.И. Аверин. – М.: Машиностроение, 1990. – 320 с.

9. R.O. Duda, P.E.Hart and D.G. Stork, "Pattern classification", Wiley, 2008, 738 p.

10. Гороховатський В.А. Структурний аналіз і інтелектуальна обробка даних в комп'ютерному зренні: монографія / В.А. Гороховатський. – Х.: Компанія СМІТ, 2014. – 316с.

11. A. Zamula and S. Kavun, Complex systems modeling with intelligent control elements, Int. J. Model. Simul. Sci. Comput. 08(01), 1750009 (2017). [Електронний ресурс]. – Режим доступу:

<https://www.worldscientific.com/doi/abs/10.1142/S179396231750009X>

12. OpenCV Open Source Computer Vision. [Электронный ресурс]. – Режим доступа: <https://docs.opencv.org/master/index.html>

(GADETSKA Svitlana) кандидат фізико-математичних наук, завідувач кафедри інформаційних технологій, Харків, Україна,
e-mail: svgadetska@ukr.net;

ПОНОМАРЕНКО Харківський національний університет радіоелектроніки, Харків, Україна,
Роман Петрович Студент-магістрант кафедри інформатики,
(PONOMARENKO e-mail: roman.ponomarenko.ubermench@gmail.com;
Roman)

Рецензент: д-р техн. наук, проф. Є.П. Путятін, Харківський національний університет радіоелектроніки, м. Харків

СТАТИСТИЧЕСКИЕ РАСПРЕДЕЛЕНИЯ И ЦЕПНОЕ ПРЕДСТАВЛЕНИЕ ДАННЫХ ПРИ ОПРЕДЕЛЕНИИ РЕЛЕВАНТНОСТИ СТРУКТУРНЫХ ОПИСАНИЙ ВИЗУАЛЬНЫХ ОБЪЕКТОВ

В.А. Гороховатский, С.В. Гадецкая, Р.П. Пономаренко

***Предметом** исследования статьи являются модели для установления уровня релевантности изображений в пространстве распределений для дескрипторов ключевых точек при распознавании визуальных объектов в системах компьютерного зрения. **Целью** является создание метода структурного распознавания изображений на основании внедрения цепных моделей данных с использованием вероятностных распределений множества дескрипторов. **Задачи:** разработка математических и программных моделей эффективного по быстродействию анализа данных при определении релевантности структурных описаний, изучение свойств, атрибутов применения, значений параметров этих моделей, оценивание эффективности по результатам обработки конкретных изображений. Применяемыми методами являются: детектор BRISK для формирования дескрипторов ключевых точек, интеллектуальный анализ данных, методы побитовой обработки и построения распределений битовых данных, аппарат метрического определения релевантности, программное моделирование. Получены следующие **результаты**. Переход от описания множеств дескрипторов к вероятностным распределениям фрагментов и сопоставление образов в пространстве распределений обеспечивают необходимую результативность распознавания. Обработка и анализ данных выполняются в сотни раз быстрее, чем традиционный подсчет голосов точек. Обработка и анализ сочетаний битов формирует значимые свойства для совокупности элементов описания с сохранением структуры данных и их унификации. С увеличением числа битов во фрагменте распределения растет расстояние между изображениями, что способствует увеличению степени их различия. Цепным представлением и применением распределений создается новое пространство данных, что позволяет существенно улучшить показатели функционирования систем распознавания изображений. **Выводы.** Научная новизна исследования заключается в усовершенствовании метода структурного распознавания изображений на основе внедрения обобщенной цепной структуры описания с использованием значений вероятностного распределения для фрагментов множества дескрипторов ключевых точек, которые содержательно отражают свойства изображений объектов и обеспечивают результативное распознавание. Практическая значимость – достижение существенного уровня повышения быстродействия вычисления релевантности, подтверждение результативности предложенных модификаций на примерах изображений, получение прикладных программных моделей для исследования и внедрения методов классификации в системах компьютерного зрения.*

***Ключевые слова:** структурные методы распознавания изображений, ключевая точка, детектор BRISK, цепное представление, распределение данных фрагмента, общий дескриптор, релевантность описаний, голосование, манхэттенская метрика, быстродействие определения релевантности.*

STATISTICAL DISTRIBUTIONS AND CHAIN REPRESENTATION OF DATA WHEN DETERMINING THE RELEVANCE OF STRUCTURAL DESCRIPTIONS OF VISUAL OBJECTS

V. O. Gorokhovatskyi, S.V. Gadetska, R.P. Ponomarenko

***The subjects** of the paper are the models for estimation of the relevance between images in the space of key point descriptors when recognizing visual objects in computer vision systems. **The goal** is to create an image structural recognition method based on the implementation of chain data models using probability distributions of the sets of descriptors. **The tasks** include the development of mathematical and software models of efficient data analysis for determining the relevance of structural descriptions, investigation of the properties, application attributes, values of parameters of these models, evaluation of the effectiveness of the specific image processing. **The methods** are used: a BRISK detector for forming the key point descriptors, data mining, methods of bitwise processing and building bit-data distributions, a method of metric relevance estimation, software modeling. The following **results** were obtained. The transition from the sets of descriptors to probability distributions of fragments and the comparison of images in the space of distributions provide the necessary recognition performance. Data processing and analysis are performed hundreds of times faster than traditional vote counting. Processing and analysis of bit combinations forms significant properties for a set of description elements with keeping the data structure and their unification. With an increase of the number of bits in the distribution fragment, the distance between images increases and it contributes to an increase of*

their difference degree. The chain representation and the use of distributions create a new data space, which allows to improve the performance of image recognition systems significantly. **Conclusions.** The contribution of the paper is the improvement of the structural image recognition method with the introduction of a generalized chain description structure using the distribution values for fragments of the set of key point descriptors, which meaningfully reflect the properties of image objects and provide effective recognition. The practical significance of the paper is the achievement of an increase of image relevance calculation speed, confirmation of the effectiveness of proposed modifications on sample images, obtaining of an application software models for research and implementation of classification methods in computer vision systems.

Keywords: structural image recognition methods, key point, BRISK detector, chain representation, fragment data distribution, general descriptor, descriptive relevance, voting, Manhattan metric, speed of relevancy estimation.