



Є.В. Бодянський¹, А.Ю. Шафроненко², І.М. Климова³

¹Доктор технічних наук, професор кафедри Штучного інтелекту.

Харківський національний університет радіоелектроніки
yevgeniy.bodyanskiy@nure.ua, ORCID 0000-0001-5418-2143

²Кандидат технічних наук, доцент кафедри Інформатики.

Харківський національний університет радіоелектроніки
alina.shafronenko@nure.ua, ORCID 0000-0002-8040-0279

³Асистент кафедри Системотехніки.

Харківський національний університет радіоелектроніки
iryna.klymova@nure.ua, ORCID 0000-0003-0455-6180

РЕКУРЕНТНА ДОСТОВІРНА НЕЧІТКА КЛАСТЕРИЗАЦІЯ ВЕЛИКИХ ДАНИХ З ВИКОРИСТАННЯМ ФУНКЦІЇ НАЛЕЖНОСТІ СПЕЦІАЛЬНОГО ТИПУ

Запропоновано метод достовірної нечіткої кластеризації для задач, коли дані надходять на обробку або у послідовному онлайн режимі, або формують надвеликі масиви (Big Data). Введені процедури є за суттю градієнтними алгоритмами оптимізації цільової функції спеціального виду, та мають низку переваг перед відомими ймовірнісними та можливісними підходами і, перш за все, робастність до аномальних спостережень. В основі підходу лежить використання міри подібності, параметри якої визначаються автоматично у процесі самонавчання. Запропоновані процедури є узагальненням відомих методів, характеризуються високою швидкістю та є досить простими у чисельній реалізації.

НЕЧІТКА КЛАСТЕРИЗАЦІЯ, ДОСТОВІРНА НЕЧІТКА КЛАСТЕРИЗАЦІЯ, РІВЕНЬ НАЛЕЖНОСТІ, РІВЕНЬ ДОСТОВІРНОСТІ, МІРА ПОДІБНОСТІ

Бодянский Е.В., Шафроненко А.Ю., Климова И.М. Рекуррентная достоверная нечеткая кластеризация больших данных с использованием функции принадлежности специального типа. Предложен метод достоверной нечеткой кластеризации для задач, когда данные поступают на обработку или в последовательном онлайн режиме, или формируют большие массивы (Big Data). Введенные процедуры являются по сути градиентными алгоритмами оптимизации целевой функции специального вида, и имеют ряд преимуществ перед известными вероятностными и возможностными подходами и, прежде всего, робастность к аномальным наблюдениям. В основе подхода лежит использование меры сходства, параметры которых определяются автоматически в процессе самообучения. Предложенные процедуры являются обобщением более известных методов, характеризуются высокой скоростью и являются достаточно простыми в численной реализации.

НЕЧЕТКАЯ КЛАСТЕРИЗАЦИЯ, ДОСТОВЕРНАЯ НЕЧЕТКАЯ КЛАСТЕРИЗАЦИЯ, УРОВЕНЬ ПРИНАДЛЕЖНОСТИ, УРОВЕНЬ ДОСТОВЕРНОСТИ, МЕРА ПОДОБИЯ

Bodyanskiy Ye.V., Shafronenko A.Yu., Klymova I. M. Recurrent authenticity of the clustering of great tribute to the function of special type. A method of credibilistic fuzzy clustering is proposed for problems when data are fed sequentially, in online mode and forms large arrays (Big Data). The introduced procedures are essentially gradient algorithms for optimizing the objective function of a special type, and have a number of advantages over known probabilistic and possible approaches and, above all, robustness to anomalous observations. The approach is based on similarity measure, parameters of that are determined automatically in the process of self-learning. The proposed procedures are a generalization of the known methods, characterized by high speed and simple in numerical implementation.

FUZZY CLUSTERING, CREDIBILISTIC FUZZY CLUSTERING, MEMBERSHIP LEVEL, CREDIBILISTIC LEVEL, SIMILARITY MEASURE

Вступ

Задача кластеризації (класифікації в режимі самонавчання) багатовимірних даних є важливою частиною інтелектуального аналізу даних (Data Mining), в рамках якої склався ряд напрямків і підходів [1, 2]. Один з таких напрямків утворюють методи нечіткої (фаззі-) кластеризації, в основі яких лежить припущення про те, що класи — кластери, що формуються взаємно перетинаються так, що кожне спостереження — вектор з різними рівнями належності — ймовірності — можливості може належати одночасно до кількох чи всіх класів.

Тут найбільш широкого поширення набули алгоритми ймовірнісної нечіткої кластеризації і, перш за все, метод нечітких С-середніх (FCM) [3,4]. Можливості цього підходу обмежуються ймовірнісними обмеженнями на рівні належності так, що «забруднені» збуреннями і викидами спостереження можуть бути віднесені до різних класів з практично однаковими рівнями належності.

У зв'язку з цим в [5] був запропонований можливісний підхід до нечіткої кластеризації (PCM) більш стійкий до шумів і збурень. Разом з тим PCM-алгоритми страждають від, так званої, проблеми

співпадіння, коли в процесі обробки інформації деякі кластери починають зливатися один з одним, що в результаті веде до невірної оцінки кількості цих кластерів.

Цих недоліків позбавлені алгоритми достовірної нечіткої кластеризації [6-8], засновані на апараті теорії вірогідності [9]. В рамках цього підходу в процесі розрахунків оцінюються як рівні нечіткої належності, так і рівні довіри, засновані на мірі належності спеціального вигляду [10]. Результати експериментів показали [7,8], що достовірний підхід забезпечує більш високу якість кластеризації в порівнянні з ймовірнісними і можливісними методами.

Вихідною інформацією для рішення задачі нечіткої кластеризації є масив n -вимірних векторів спостережень $X = \{x_1, x_2, \dots, x_N\} \subset R^n$, $x(k) \in X$, $k=1, 2, \dots, N$, який повинен бути розбитий на m кластерів-кластерів з деяким рівнем належності-можливості-достовірності $U_q(k)$ k -го вектора x_k до q -го кластеру ($1 < m < N, 1 \leq q \leq m$). Необхідно також відзначити, що вихідні дані попередньо предоброблені так, що $-1 \leq x_{ki} \leq 1$ ($1 \leq i \leq n$), де x_{ki} — i -та компонента вектора x_{ki} .

Таким чином, задача кластеризації вирішується в пакетному режимі, коли весь масив даних обробляється багаторазово на основі почергового оцінювання [8].

Якщо ж дані надходять на обробку у вигляді потоку або утворюють надвеликі масиви, що досить часто зустрічаються у реальних ситуаціях, пакетний режим не дозволяє ефективно вирішити розглянуту задачу.

У цій ситуації найбільш ефективними є рекурентні процедури нечіткої кластеризації, що дозволяють вирішувати завдання в online режимі і уточнювати отримані рішення по мірі надходження кожного нового спостереження. Так, у [9, 10] були запропоновані рекурентні варіанти FCM, які є за суттю градієнтними процедурами оптимізації прийнятої цільової функції, а в [11, 12] були введені рекурентні модифікації PCM, що призначені для послідовної обробки даних.

У зв'язку з цим є доцільною розробка рекурентної модифікації методу достовірної нечіткої кластеризації, що дозволяє уточнювати шукані характеристики кластерів по мірі надходження кожного нового спостереження.

1. Рекурентний алгоритм достовірної нечіткої кластеризації

Найбільш популярний метод ймовірнісної нечіткої кластеризації пов'язаний з мінімізацією цільової функції [4]

$$E(U_q(k), w_q) = \sum_{k=1}^N \sum_{q=1}^m U_q^\beta(k) D^2(x_k, w_q) \quad (1)$$

за обмежень

$$\sum_{q=1}^m U_q(k) = 1, \\ 0 < \sum_{q=1}^m U_q(k) < N.$$

Вирішуючи задачу нелінійного програмування за допомогою методу невизначених множників Лагранжа, приходимо до відомого результату

$$U_q^{(\tau+1)}(k) = (D^2(x_k, w_q^{(\tau)}))^{\frac{1}{1-\beta}} \times \\ \times \left(\sum_{l=1}^m (D^2(x_k, w_l^{(\tau)}))^{\frac{1}{1-\beta}} \right)^{-1}, \quad (2)$$

$$w_q^{(\tau+1)} = \sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^\beta x_k \left(\sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^\beta \right)^{-1}, \quad (3)$$

де $U_q(k)$ — рівень нечіткої належності спостереження x_k до q -го кластера; Cl_q ($1 \leq q \leq m$), w_q — прототип-центроїд q -го кластера; $\beta > 1$ — параметр фазифікації, що задає «розмитість» границь кластерів; $D(x_k, w_q)$ — відстань між x_k та w_q у прийнятій метриці; $\tau = 0, 1, 2, \dots$ — індекс епохи обробки інформації в режимі почергового оцінювання.

При цьому процес обчислень триває до виконання умови $w_q^{(\tau+1)} - w_q^{(\tau)} \leq \varepsilon \forall 1 \leq q \leq m$, де ε — наперед заданий поріг точності обчислень.

У випадку $\beta=2$ та евклідової метрики

$$D^2(x_k, w_q) = \|x_k - w_q\|_2^2$$

приходимо до популярного алгоритму нечітких C-середніх (FCM) [13] вигляду

$$U_q^{(\tau+1)}(k) = \|x_k - w_q^{(\tau)}\|^{-2} \left(\sum_{l=1}^m \|x_k - w_l^{(\tau)}\|^{-2} \right)^{-1}, \quad (4)$$

$$w_q^{(\tau+1)} = \sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^2 x_k \left(\sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^2 \right)^{-1}. \quad (5)$$

Якщо дані надходять на обробку послідовно в онлайн режимі, задача нелінійного програмування може бути вирішена за допомогою алгоритму Ерроу-Гурвіца-Удзави, що є за суттю градієнтною процедурою пошуку сідлової точки функції Лагранжа на основі критерію (1) з обмеженнями на суму належностей.

При цьому співвідношення (2), (3) можуть бути переписані у формі

$$\begin{cases} U_q(k+1) = (D^2(x_{k+1}, w_q^{(k)}))^{\frac{1}{1-\beta}} * \\ * \left(\sum_{l=1}^m (x_{k+1}, w_l^{(k)}) \right)^{\frac{1}{1-\beta}})^{-1}, \\ w_q(k+1) = w(k) + \eta(k+1) U_q^\beta(k+1) (x_{k+1} - w_q(k)), \end{cases} \quad (6)$$

де $\eta(k)$ — параметр кроку навчання, а (4), (5) —

$$\begin{cases} U_q(k+1) = \|x_k - w_q(k)\|^{-2} \times \\ \times \left(\sum_{l=1}^m \|x_k - w_l(k)\|^{-2} \right)^{-1}, \\ w_q(k+1) = w_q(k) + \eta(k+1) U_q^2(k+1) \times \\ \times (x_{k+1} - w_q(k)), \end{cases} \quad (7)$$

що є узагальненням рекурентних процедур Парка-Деггера [9] і Чанга-Лі [10].

Можливісні алгоритми нечіткої кластеризації засновані на мінімізації цільової функції [5]

$$\begin{aligned} E(U_q(k), w_q, \mu_q) = & \sum_{k=1}^N \sum_{q=1}^m U_q^\beta(k) D^2(x_k, w_q) + \\ & + \sum_{q=1}^m \mu_q \sum_{k=1}^N (1 - U_q(k))^\beta, \end{aligned} \quad (8)$$

де $\mu_q \geq 0$ визначає відстань, на якій рівень належності приймає значення 0,5, тобто $U_q(k) = 0$, якщо $D^2(x_k, w_q) = \mu$.

Мінімізація критерію (8) дозволяє отримати аналітичний розв'язок у вигляді

$$U_q^{(\tau+1)}(k) = \left(1 + \left(\frac{D^2(x_k, w_q^{(\tau)})}{\mu_q^{(\tau)}} \right)^{\frac{1}{\beta-1}} \right)^{-1}, \quad (9)$$

$$w_q^{(\tau+1)} = \sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^\beta x_k \left(\sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^\beta \right)^{-1}, \quad (10)$$

$$\begin{aligned} \mu_q^{(\tau+1)} = & \sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^\beta D^2(x_k, w_q^{(\tau+1)}) \times \\ & \times \left(\sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^\beta \right)^{-1}, \end{aligned} \quad (11)$$

яке у квадратичному випадку набуває вигляду

$$U_q^{(\tau+1)}(k) = \left(1 + \frac{\|x_k - w_q^{(\tau)}\|^2}{\mu_q^{(\tau)}} \right)^{-1}, \quad (12)$$

$$w_q^{(\tau+1)} = \sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^2 x_k \left(\sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^2 \right)^{-1}, \quad (13)$$

$$\mu_q^{(\tau+1)} = \sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^2 \|x_k - w_q^{(\tau+1)}\|^2 \left(\sum_{k=1}^N \left(U_q^{(\tau+1)}(k) \right)^2 \right)^{-1}. \quad (14)$$

Онлайн версії (9)-(14) при цьому мають вигляд [11, 12]

$$\begin{cases} U_q(k+1) = \left(1 + \left(\frac{D^2(x_{k+1}, w_q(k))}{\mu_q(k)} \right)^{\frac{1}{\beta-1}} \right)^{-1}, \\ w_q(k+1) = w_q(k) + \eta(k+1) U_q^\beta(k+1) * (x_{k+1} - w_q(k)), \\ \mu_q(k+1) = \sum_{p=1}^{k+1} U_q^\beta(p) D^2(x_p, w_q(k+1)) * \left(\sum_{p=1}^{k+1} U_q^\beta(p) \right)^{-1} \end{cases} \quad (15)$$

та (при $\beta = 2$)

$$\begin{cases} U_q(k+1) = \left(1 + \frac{\|x_{k+1} - w_q(k)\|^2}{\mu_q(k)} \right)^{-1}, \\ w_q(k+1) = w_q(k) + \eta(k+1) U_q^2(k+1) * (x_{k+1} - w_q(k)), \\ \mu_q(k+1) = \sum_{p=1}^{k+1} U_q^2(p) \|x_p - w_q(k+1)\|^2 * \left(\sum_{p=1}^{k+1} U_q^2(p) \right)^{-1}. \end{cases} \quad (16)$$

Достовірні нечіткі кластеризації пов'язані з мінімізацією цільової функції

$$E(Cr_q(k), w_q) = \sum_{k=1}^N \sum_{q=1}^m Cr_q^\beta(k) D^2(x_k, w_q) \quad (17)$$

за обмежень

$$0 \leq Cr_q(k) \leq 1 \forall q, k;$$

$$\sup Cr_q(k) \geq 0,5 \forall k;$$

$$Cr_q(k) + \sup Cr_r(k) = 1$$

для будь яких q та k , для яких $Cr_q(k) \geq 0,5$. Тут $Cr_q(k)$ — рівень нечіткої достовірності того, що спостереження x_k належить кластеру U_q . При цьому рівень достовірності розраховується на основі функції належності [14]

$$U_q(k) = \gamma_q(D(x_k, w_q)) \quad (18)$$

що задовольняє умовам:

- $\gamma_q(\cdot)$ монотонно зменшується на інтервалі $[0, \infty]$,
- $\gamma_q(0) = 1$,
- $\gamma_q(\infty) \rightarrow 0$.

Нескладно помітити, що функція (18) є за суттю мірою подібності, заснованій на відстані [15].

В якості такої функції у [14] було запропоновано використовувати вираз

$$U_q(k) = \left(1 + D^2(x_k, w_q) \right)^{-1}, \quad (19)$$

що описує за суттю звичайну дзвонувату функцію належності, яка використовується в системах нечіткого висновування.

Цікаво зауважити, що вираз (2) може бути переписано у формі

$$\begin{aligned} U_q(k) = & \left(D^2(x_k, w_q(k)) \right)^{\frac{1}{1-\beta}} \times \\ & \times \left(\sum_{l=1}^m \left(D^2(x_k, w_l(k)) \right)^{\frac{1}{1-\beta}} \right)^{-1} = \\ = & \left(D^2(x_k, w_q(k)) \right)^{\frac{1}{1-\beta}} \left(D^2(x_k, w_q(k)) \right)^{\frac{1}{1-\beta}} + \\ & + \sum_{l=1, l \neq q}^m \left(D^2(x_k, w_l(k)) \right)^{\frac{1}{1-\beta}})^{-1} = \\ = & \left(1 + \left(D^2(x_k, w_q(k)) \right)^{\frac{1}{1-\beta}} \sum_{l=1, l \neq q}^m \left(D^2(x_k, w_l(k)) \right)^{\frac{1}{1-\beta}} \right)^{-1}, \end{aligned} \quad (20)$$

яка для евклідової метрики $\beta = 2$ приймає вигляд функції щільності розподілу Коші з параметром ширини σ_q^2 [16]:

$$U_q(k) = \left(1 + \frac{\|x_k - w_q(k)\|^2}{\sigma_q^2} \right)^{-1}, \quad (21)$$

$$\sigma_q^2 = \left(\sum_{\substack{l=1 \\ l \neq q}}^m \|x_k - w_l(k)\|^{-2} \right)^{-1}. \quad (22)$$

Нескладно бачити, що функція належності (19) є окремим випадком (21) при $\sigma_q^2 = 1$.

Остаточний пакетний алгоритм достовірної нечіткої кластеризації може бути записаний у вигляді [7,8]:

$$U_q^{(\tau+1)}(k) = \left(1 + D^2(x_k, w_q^{(\tau)}) \right)^{-1}, \quad (23)$$

$$U_q^{*(\tau+1)}(k) = U_q^{(\tau+1)}(k) \left(\sup_{l \neq q} U_l^{(\tau+1)}(k) \right)^{-1}, \quad (24)$$

$$Cr_q^{(\tau+1)}(k) = \frac{1}{2} \left(U_q^{*(\tau+1)}(k) + 1 - \sup_{l \neq q} U_l^{*(\tau+1)}(k) \right), \quad (25)$$

$$w_q^{(\tau+1)} = \sum_{k=1}^N \left(Cr_q^{(\tau+1)}(k) \right)^\beta x_k \left(\sum_{k=1}^N \left(Cr_q^{(\tau+1)}(k) \right)^\beta \right)^{-1}. \quad (26)$$

На підставі (17), (21) — (26) можна ввести у розгляд онлайн рекурентну версію алгоритму достовірної нечіткої кластеризації у вигляді

$$\begin{cases} \sigma_q^2(k+1) = \left(\sum_{\substack{l=1 \\ l \neq q}}^m \|x_{k+1} - w_l(k)\|^{-2} \right)^{-1}, \\ U_q(k+1) = \left(1 + \frac{\|x_{k+1} - w_q(k)\|^2}{\sigma_q^2(k+1)} \right)^{-1}, \\ U_{(k+1)}^* = U_q(k+1) \left(\sup_{l \neq q} U_l(k+1) \right)^{-1}, \\ Cr_q(k+1) = \frac{1}{2} \left(U_q^*(k+1) + 1 - \sup_{l \neq q} U_l^*(k+1) \right), \\ w_q(k+1) = w_q(k) + \eta(k+1) Cr_q^\beta(k+1) * \\ * (x_{k+1} - w_q(k)). \end{cases} \quad (27)$$

Як можна бачити, з обчислювальної точки зору online алгоритм достовірної нечіткої кластеризації не складніший рекурентних версій FCM і РСМ, зберігаючи при цьому переваги достовірного підходу.

2. Результати обчислювального експерименту

Щоб перевірити ефективність, оцінити працездатність та спроможність якісно кластеризувати велику кількість даних розробленого методу, а також довести його перевагу над аналогами, було проведено експериментальне дослідження за допомогою двох різних баз даних.

Проведено порівняльний аналіз якості кластеризації даних за основними характеристиками

рейтингів якості таких, як: коефіцієнт розподілу (РС) який визначає «перекриття» між кластерами; індекс розподілу (SC), що кількісно визначає співвідношення суми компактності та розділення кластерів; індекс Ксі — Бені (XB) визначає співвідношення загальної мінливості всередині кластерів та їх поділу до існуючих методів кластеризації та запропонованого методу.

Для порівняння було обрано найбільш відомі методи кластеризації, такі як: нечіткі С-середні (FCM), алгоритм Густафсона-Кесселя (GK), алгоритм Гата-Геви (GG), метод адаптивної ймовірнісної нечіткої кластеризації, метод адаптивної нечіткої можливісної кластеризації даних та метод адаптивної нечіткої кластеризації даних.

Результати експериментальних досліджень та порівняльного аналізу із зазначеними характеристиками наведені в табл. 1 та табл. 2, для першої та другої вибірок відповідно.

Таблиця 1

Порівняльна оцінка якості кластеризації нечітких методів кластеризації з використанням першого набору даних

Методи кластеризації даних	Перша вибірка		
	PC	SC	XB
FCM	0,50	1,62	0,19
Густафсон-Кессель	0,27	1,66	1,62
Гат-Гева	0,25	1,54	1,35
Адаптивна ймовірнісна нечітка кластеризація	0,25	1,44	0,01
Адаптивна нечітка можливісна кластеризація	0,26	1,22	0,01
Рекурентна нечітка достовірна кластеризація	0,21	1,13	0,01

Таблиця 2

Порівняльна оцінка якості кластеризації нечітких методів кластеризації з використанням другого набору даних

Методи кластеризації даних	Друга вибірка		
	PC	PC	PC
FCM	0,48	1,60	0,19
Густафсон-Кессель	0,26	1,64	1,62
Гат-Гева	0,26	1,50	1,35
Адаптивна ймовірнісна нечітка кластеризація	0,25	1,42	0,01
Рекурентна нечітка можливісна кластеризація	0,37	1,11	0,18
Адаптивна нечітка достовірна кластеризація	0,23	1,22	0,01

Аналізуючи та оцінюючи отримані результати експериментальних досліджень, можна зробити висновки, що запропонований метод рекурентної достовірної нечіткої кластеризації великих даних з використанням функції належності спеціального типу має достатньо якісні результати кластеризації, що підтверджується експериментально.

Увагу привертає ще те, що запропонований підхід достатньо простий в реалізації та практично по всіх критеріях аналізу кластерування даних, не поступається більш відомим нечітким алгоритмам.

На рис. 1 та 2 більш наочно продемонстрована робота запропонованого методу рекурентної достовірної нечіткої кластеризації великих даних з використанням функції належності спеціального типу.

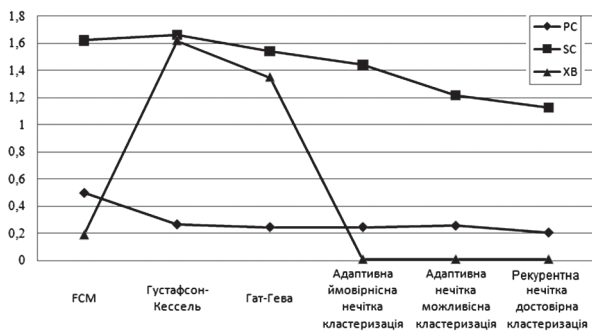


Рис. 1. Порівняльний аналіз роботи методів кластеризації для першої вибірки

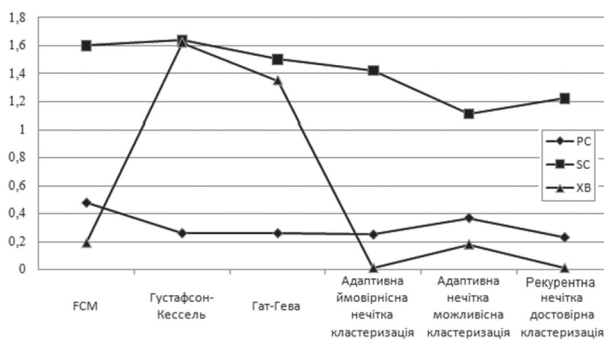


Рис. 2. Порівняльний аналіз роботи методів кластеризації для другої вибірки

Висновки

Розглянуто задачу нечіткої кластеризації на основі ймовірнісного, можливісного і достовірного підходів на основі пакетного і online режимів надходження і обробки інформації. Введена рекурентна версія достовірного алгоритму, що є за суттю процедурою градієнтної оптимізації прийнятого критерія нечіткої достовірної кластеризації. Введена модифікація функції належності, що є по суті мірою подібності та узагальненням відомих раніше функцій. Запропонований метод рекурентної достовірної нечіткої кластеризації великих даних з використанням функції належності спеціального типу є доволі простим в чисельній реалізації і призначений для вирішення завдань, що виникають у рамках інтелектуального аналізу великих даних.

Список літератури:

- [1] Xu R., Wunsch D.C. II. Clustering— Hoboken, N.J.: John Wiley & Sons, Inc., 2009.
- [2] Aggarwal C.C. Data Mining: Text Book. Springer, 2015.
- [3] Bezdek J.C. Pattern Recognition with Fuzzy Objective Function Algorithms.-N.Y.:Plenum Press, 1981.

- [4] Höppner F., Klawonn F., Kruse R., Runkler T. Fuzzy Clustering Analysis: Methods for Classification, Data Analysis and Image Recognition.-Chichester: John Wiley & Sons, 1999.-289p.
- [5] R. Krishnapuram, J.M. Keller. A possibilistic approach to clustering. Fuzzy Systems, 1993, 1, №2, P.98-110.
- [6] Chintalapudi K. K. and M. Kam, "A noise resistant fuzzy c-means algorithm for clustering," IEEE Conference on Fuzzy Systems Proceedings, vol. 2, May 1998, pp. 1458-1463.
- [7] Zhou J., Wang Q., Hung C.-C., Yi X. Credibilistic clustering: the model and algorithms. Int.J. of Uncertainty, Fuzziness and Knowledge-Based Systems- 2015-23-№4- P.545-564.
- [8] Zhou, J., Wang, Q., Hung, C. C. Credibilistic clustering algorithms via alternating cluster estimation.- J. Intell. Manuf.-2017-28-P.727-738.
- [9] Liu, B., & Liu, Y. Expected value of fuzzy variable and fuzzy expected value models. IEEE Transactions on Fuzzy Systems, — 2002-10-№ 4-P. 445-450.
- [10] Liu, B. A survey of credibility theory. Fuzzy Optimization and Decision Making-2006-5-№4-P. 387-408.
- [11] D.C. Park, I. Dagher. Gradient based fuzzy c-means (GBFCM) algorithm. Proc. IEEE Int. Conf. on Neural Networks, 1984, P.1626-1631.
- [12] F.L. Chung, T. Lee. Fuzzy competitive learning. Neural Networks, 1994, 7, №3, P.539-552.
- [13] Bodyanskiy Ye, Kolodyazhnyi V., Stephan A. Recursive fuzzy clustering algorithms. —Proc 10th East West Fuzzy Coll. 2002, -Zittau- Görlitz, HS, 2002-P.276-283.
- [14] Bodyanskiy, Ye. Computational intelligence techniques for data analysis / Ye. Bodyanskiy // Lecture Notes in Informatics.-Bonn: V-P- 72, GI, 2005. — P.15-36.
- [15] Zhou, J., & Hung, C.-C. (2007). A generalized approach to possibilistic clustering algorithms. Int. J. of Uncertainty, Fuzziness and Knowledge-Based Systems. — 2007 — 15. — P. 117-138.
- [16] Young F.W., Hamer R.M. Theory and Applications of Multidimensional Scaling-Hillsdale, N.J.: Erlbaum, 1994.
- [17] Hu Zh., Bodyanskiy Ye, Tyshchenko O., Shafronenko A. Fuzzy clustering of incomplete data by means of similarity measures- Proc.2019 IEEE 2nd Ukr. Conf. on Electrical and Computer Engineering (UKRCON), Track 6.-Lviv, Ukraine, 2019.-P.149-152.
- [18] Bezdek J.C. A convergence theorem for the fuzzy ISODATA clustering algorithms. — IEEE Trans. Pattern Anal. Mach. Intell. — 1980 — 2. — P. 1- 8.
- [19] Bodyanskiy Ye., Gorshkov Ye., Kokshenev I., Kolodyazhnyi V. Outlier resistant recursive fuzzy clustering algorithms. Ed. by B. Reusch «Computational Intelligence Theory and Applications» — Advances in Soft Computing-Vol.38.-Berlin Heidelberg, Springer Verlag, 2006-P.647-652.
- [20] Bodyanskiy Ye., Gorshkov Ye., Kokshenev I., Kolodyazhnyi V. Robust recursive fuzzy clustering algorithms- Proc. 12th East West Fuzzy Coll 2005 — Zittau- Görlitz, FH, 2005-P.301-308.
- [21] Bodyanskiy Ye, Shafronenko A., Mashtalir S., Online robust fuzzy clustering of data with omissions using similarity measure of special type — Lecture Notes in Computational Intelligence and Decision Making-Cham: Springer, 2020-P.637-646.

Надійшла до редколегії 18.11.2020