

Висвітлено проблему формалізації об'єкта розпізнавання, що пов'язані з накладанням об'єктів, перетином інтенсивностей пікселів і фону, неоднорідності фону, розмитості границь об'єкта. Остання проблема впливає на задачу класифікації, оскільки залежно навіть від малої кількості (зайвих або нестачі) пікселів приналежність до класу може змінитись.

1. Грицик В.В., Влах М.А. Технічні та програмні засоби розпізнавання та аналізу зображень складних біологічних об'єктів. – Львів.: ІТІС. 2005, Т.8. – №1. – С.17–28. 2. Фотографії та методика аналізу для побудови алгоритму та системи автоматизації експерименту надані чл.-кор. НАН України д.б.н. Р.С. Стойкою та працівником інституту біології клітини (м.Львів) О.Ю. Ключівською (2006 рік). 3. www.probes.com; пропозиції провідних виробників програмного забезпечення, в т. ч. фірма Molecular Probes. 4. www.videotest.ru/ru

УДК 004.93'14

Р. Мельник, Р. Тушницький

Національний університет “Львівська політехніка”
кафедра програмного забезпечення

КАСКАДНА ДЕКОМПОЗИЦІЯ МНОЖИН ВЕЛИКОЇ РОЗМІРНОСТІ ПРИ КЛАСТЕРИЗАЦІЇ КЛЮЧІВ ОБРАЗІВ

© Мельник Р., Тушницький Р., 2007

Розглянуто застосування методу кластеризації до розв'язання задачі зберігання та пошуку ключів, що відображають графічні образи. З метою зменшення алгоритмічної складності програми формування кластерів візуальних образів розроблено підхід каскадної декомпозиції множин ключів та їх кластерів. Продемонстровано результати дослідження впливу коефіцієнтів ділення на якісні показники розбиття.

The clustering algorithm for storage and searching of visual patterns keys is considered. To reduce algorithmic complexity of the application to formulate the visual patterns clusters the cascade decomposition approach for keys and clusters is proposed. Some results with quality characteristics of decomposition depending from partitioning parameters are presented.

Вступ. Задачі розпізнавання образів вимагають значних часових затрат через свою розмірність та складність описів образів. Ключі образів мають менше ознак, але через їх велику кількість задача їх об'єднання та пошук залишається складною з погляду часу. Для зменшення алгоритмічної складності в статті пропонується підхід багатокаскадної декомпозиції множин ключів та їх кластерів на певних рівнях ієрархічного дерева згортання кластерів. Кількість рівнів та потужності множин для учасників рекурсивного згортання є параметрами управління процесом декомпозиції ключів.

Постановка задачі кластеризації ключів образів. Методи кластерного аналізу широко використовуються для декомпозиції, дослідження та розпізнавання зображень [1 – 5]. Зокрема, робота [1] містить класифікацію методів кластеризації та спосіб формування контурів виділених кластерів. Робота [2] присвячена кластеризації графових моделей, якими відображають частини зображень. В роботі [3] до зображень застосовано підхід кластеризації експериментальних даних, на які накладено сітку з певним кроком. У наведених роботах, як і в інших, немає розв'язку задач

кластеризації великих обсягів експериментальних даних з метою зменшення часових затрат та функцій оцінки показників отриманого розбиття як способів керування ними.

Образи, призначені до зберігання та розпізнавання, представимо множиною $Q(Q_1, Q_2, Q_3, \dots, Q_N)$, де N – достатньо велике значення. Кожному образу притаманні певні характеристичні ознаки $a, b, c, \dots$ (координати центра, кількості ліній перетину ліній образу з осями координат тощо), які формують ключ образу і виражаються числовими значеннями. Тобто кожному образу відповідає множина його характеристик $Q_1 = Q_1(a_1, b_1, c_1, \dots)$, $Q_2 = Q_2(a_2, b_2, c_2, \dots)$, $Q_3 = Q_3(a_3, b_3, c_3, \dots)$, $\dots$, $Q_N = Q_N(a_N, b_N, c_N, \dots)$. Для прискорення пошуку відповідного образу (групи образів) у списку образів останні згрупуємо, зафіксувавши місце розташування групи на дереві. Групи образів відповідають кластерам, утвореним за певними критеріям близькості на основі дерева згортання зображення [2, 3].

Кластеризація здійснюється процедурою, на i -му кроці якої формується множина X_{i+1} вершин дерева (крім початкових образів), що представляється так:

$$(\forall Q_k, Q_j \in X_i) [F^*(Q_k, Q_j) = \text{opt } F(X')], X' \in S, \\ (\forall Q_e = Q_k \cup Q_j \in B_i), \quad (1)$$

$C_i = X_i \setminus Q_e$, тобто множина B_i утворюється на основі новоутворених вершин, які задовольняють найкращі значення функції критерію $F(X')$. Крім того, до цієї множини входять вершини множини C_i попереднього рівня, які не об'єднувались, тобто

$$X_{i+1} = B_i \cup C_i. \quad (2)$$

Об'єднання двох ключів образів відбувається, якщо функція для них набуває мінімального значення

$$F^* = \min(F_{kj}), k, j \in I, \quad (3)$$

де I – множина всіх можливих пар ключів початкових образів, кластерів та їх комбінацій. Зауважимо, що новоутворені кластери образів (ключів) позначаємо також символом Q_{N+k} з індексом, більшим за кількість початкових образів N .

Прискорення алгоритму полягає в тому, щоб на кожному кроці об'єднувались ті пари ключів, критерії яких задовольняють умову:

$$F_0^*(1 + k_v) \geq F_{ij} \geq F_0,$$

де F_0 – найкраще значення функції критерію на даному кроці об'єднання, k_v ($k_v < 1$) – коефіцієнт кількості об'єднань на кроці або коефіцієнт швидкості кластеризації.

Функцію F утворимо як зважену суму модулів різниць між характеристиками образів (кластерів), що є кандидатами на об'єднання в один спільний кластер:

$$F_{ij} = w_1 |a_i - a_j| + w_2 |b_i - b_j| + w_3 |c_i - c_j| + \dots,$$

або як зважену суму квадратів різниць між характеристиками:

$$F_{ij} = w_1 [a_i - a_j]^2 + w_2 [b_i - b_j]^2 + w_3 [c_i - c_j]^2 + \dots,$$

де підсумовування відбувається за всіма ознаками, які формують ключ образу.

Введенням зваженої функції критерію підкреслюємо факт, що задача поділу ключів у кластери розв'язується як задача багатокритеріальної оптимізації методом згортання декількох критеріальних функцій в одну.

Важливим питанням є масштаб значень ознак ключів $a_i, b_i, c_i, \dots$, на основі яких формуються кластери. Без врахування масштабу чисел значення менших порядків нівелюються найбільшими числами. Тому кожне з значень ознак має нормалізуватись розкидами можливих значень, наперед заданих. Вважатимемо, що значення ознак ключів нормалізовані.

В процесі побудови дерева утворюються характеристики новоутворених вершин за простим усередненням характеристик складових вершин:

$$Q_k = Q_k((a_i + a_j)/2, (b_i + b_j)/2, (c_i + c_j)/2, \dots),$$

або за зваженим усередненням характеристик складових вершин для врахування розмірності кластера – претендента на об'єднання:

$$Q_k = Q_k((k \cdot a_i + r \cdot a_j)/(k+r), (k \cdot b_i + r \cdot b_j)/(k+r), (k \cdot c_i + r \cdot c_j)/(k+r), \dots),$$

де k, r – кількості ключів у i -му, j -му кластерах – вузлах дерева згортання.

Якщо не передбачити критерію зупинки, то в результаті роботи алгоритму буде побудоване повне дерево з однією вершиною – коренем. Формування кластерів ключів, що мають найближчі характеристичні ознаки, відповідає дереву неповного згортання. Побудуємо його з використанням алгоритму згортання та виділення груп на основі градієнта функції критерія об'єднання кластерів.

Для оцінювання кластерів сформулюємо такі ознаки: розміри, координати, питома густина ключа, його питомий об'єм, дисперсія. Зокрема, діаметр D_k кластера – це віддаль між максимально віддаленими ключами k -го кластера:

$$D_k = \max \{ w_1 |a_i - a_j| + w_2 |b_i - b_j| + w_3 |c_i - c_j| + \dots \}, \quad i, j \in J,$$

де J – множина всіх можливих пар номерів ключів образів у k -му кластері.

Координати центру кластера визначимо за формулою:

$$C_k = C_k (1/2 \sum a_k, 1/2 \sum b_k, 1/2 \sum c_k, \dots).$$

Інтегральною оцінкою отриманих кластерів сформулюємо як середню питому густину всіх ключів:

$$S = \{ \sum L_k / D_k \} / M,$$

або її оберненою величиною питомого об'єму на один ключ:

$$V = \{ \sum D_k / L_k \} / M.$$

де L_k, D_k – кількість ключів та діаметр k -го кластера, M – кількість отриманих кластерів.

Близьким до вказаних оцінок є функція питомої дисперсії кластерів:

$$E = \{ \sum E_k \} / M = \{ \sum \sum d_j \} / M = \{ \sum [\sum (a_i^* - a_j)^2 + (b_i^* - b_j)^2 + (c_i^* - c_j)^2 + \dots] / m_j \} / M,$$

$$k \in \overline{1, M}, \quad i, j \in J,$$

де a_i^*, b_i^*, c_i^* – координати центру простору кластера, E_k – дисперсія k -го кластера, d_j – середньоквадратичне відхилення ключа від центру простору кластера, m_j – кількість ключів у j -му кластері.

У багатовимірному просторі з різними центрами та однаковими радіусами (як частинний випадок) критерієм якості отриманих кластерів вважаємо різницю між максимальним та мінімальними діаметрами отриманих кластерів, тобто

$$R = \max D_k - \min D_k.$$

Каскадна кластеризація ключів. Обчислювальна складність розробленого алгоритму перевищує величину $O(n^2)$, що підтвердили експерименти в роботі [4] з дослідження залежності часу побудови дерева згортання від величини початкової вибірки. Для значних обсягів вхідних даних використання алгоритму пов'язане зі значними часовими затратами. Тому в роботі [5] реалізовано наближений підхід кластеризації ключів з розбиттям множини базових ключів. У роботі пропонується багаторівнева декомпозиція множин ключів і кластерів, яку назвемо каскадною кластеризацією.

Розбиваємо вхідну множину ключів $Q(Q_1, Q_2, Q_3, \dots, Q_N)$ на p підмножин $O_1(Q_1, Q_2, Q_3, \dots, Q_z), O_2(Q_{z+1}, Q_{z+2}, Q_{z+3}, \dots, Q_t), \dots, O_p(Q_{t+1}, Q_{t+2}, Q_{t+3}, \dots, Q_N)$. До кожної з підмножин (назвемо їх множинами нульового каскаду, рис. 1) застосуємо алгоритм кластеризації (1–3), утворивши множини відповідних кластерів $K_1(k_1, k_2, k_3, \dots), K_2(k_s, k_{s+1}, k_{s+2}, \dots), \dots, K_p(k_r, k_{r+1}, k_{r+2}, \dots)$, де $k_1, k_2, \dots, k_i, \dots$ – кластери, ключі в яких відносяться до відповідних підмножин $O_1, O_2, \dots, O_p$. Утворимо множину кластерів 1-го каскаду кластеризації K об'єднанням:

$$K = K_1(k_1, k_2, k_3, \dots) \cup K_2(k_s, k_{s+1}, k_{s+2}, \dots) \cup \dots \cup K_p(k_r, k_{r+1}, k_{r+2}, \dots).$$

Застосуємо до цієї множини алгоритм кластеризації (1–3), розглядаючи кожен з кластерів $k_1, k_2, \dots, k_i, \dots$ як базовий, тобто листок дерева згортання. При цьому вони (кластери 1-го каскаду) перемішуються, тобто є незалежними. Утворену множину кластерів 1-го каскаду поділимо на підмножини $O_1, O_2, \dots, O_p$, потім здійснюємо кластеризацію в межах підмножин до отримання множини кластерів 2-го каскаду. При отриманні кореня дерева згортання на 2-му каскаді матимемо

двокаскадну декомпозицію, на третьому каскаді – трикаскадну тощо. Кількість каскадів визначається кількістю рівнів ієрархічного дерева згортання, на яких здійснюється поділ кластерів на множини (без рівня ключів). Схема поділу та згортання є рекурсивною (рис.1).

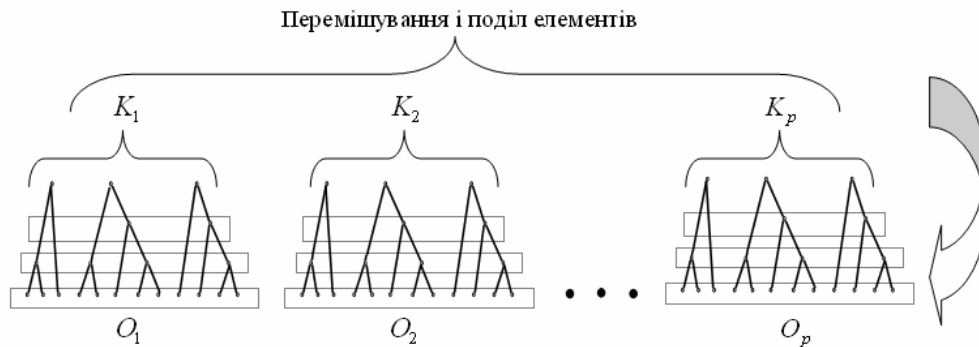


Рис.1. Каскадне дерево формування кластерів ключів

На кількість каскадів вказують коефіцієнти поділу: наприклад, 20000 ключів розбивається на 100 груп по 200 ключів з виходом на 2-й каскад по 20 кластерів з групи. Тоді на 2-му каскаді є $100 \times 20 = 2000$ кластерів, які ділимо ще на 10 груп по 200 елементів з виходом на вищий каскад 20 кластерів з групи. Кількість каскадів визначається автоматично залежно від кількості кластерів, до яких застосовується повний перебір. У розробленому пакеті це число вноситься користувачем. Питання вибору потужностей множин $O_1, O_2, \dots, O_p$ в межах каскадів i , відповідно, їх кількості залишається відкритим для дослідження.

Програмна реалізація. Для проведення досліджень складності та правильності прийнятих гіпотез побудований програмний пакет. Інтерфейс пакета дає змогу керувати процесом побудови дерева згортання. До складу пакета входять: 1) модуль керування роботою модулів побудови дерева згортання, початкового розбиття та оптимізації; 2) модуль введення/виведення та ініціалізації для забезпечення інтерпретації вхідних даних програми, побудови базових масивів вхідної інформації, виведення результатів роботи програми у внутрішньому форматі даних, збереження часової та статистичної інформації про хід виконання процесу згортання, розбиття та оптимізації; 3) модуль інтерфейсу для забезпечення діалогу з користувачем, виведення допоміжної та статистичної інформації про хід роботи та взаємозв'язку з операційною системою; 4) модуль побудови дерева згортання для генерації послідовності усіх можливих пар елементів, поновлення інформації у масиві пар та перерахунку критеріїв, модуль керує процесом об'єднання пар елементів у кластери; модуль візуалізації дерева згортання для візуального відображення дерева згортання, отримання інформації про його рівні та кластери, дослідження шляхів між кластерами, його елементами та вершиною дерева.

Аналіз результатів. При дослідженні алгоритму 10000 ключів генерувались за рівномірним законом розподілу з різними центрами координат та однаковими дисперсіями. Тестування проведені для різної кількості множин (відповідно, їх потужностей) для 0-го каскаду кластеризації. Як і очікувалось, результати, отримані на кінцевому каскаді, можна охарактеризувати так:

1) якісні та кількісні показники кластерів є тим кращими, чим більша кількість кінцевих кластерів виходить з групи, а груп є мінімальна кількість; 2) часові показники програми зворотні, тобто вони менші для меншої кількості елементів в кожній групі, а кількість груп є максимальною (у таблиці: L – кількість рівнів дерева згортання, C – кількість кластерів з групи, E – кількість елементів в групі, G – кількість груп, T – час виконання, c , CC – кількість каскадів, F – значення функції, DIS – дисперсія, CAP – питомий об'єм, DEN – питома густина, WIN – часовий виграш відносно декомпозиції без каскадів та груп, %).

Результати дослідження каскадної декомпозиції

<i>G</i>	<i>E</i>	<i>C</i>	<i>L</i>	<i>T</i> , c	<i>CC</i>	<i>F</i>	<i>DIS</i>	<i>CAP</i>	<i>DEN</i>	<i>WIN</i> , %
500	20	1	62	0,561	–	60	7,730	1,418	3,626	99,68
			39	0,400	3	149	8,107	1,418	3,044	99,78
200	50	5	82	2,103	–	25	3,338	1,141	3,001	98,82
			65	0,711	3	29	3,600	1,145	2,399	99,60
100	100	10	85	3,134	–	30	2,051	0,995	3,177	98,24
			86	1,282	2	28	2,174	0,997	2,604	99,28
50	200	20	87	4,006	–	31	1,321	0,878	3,188	97,75
			75	2,494	2	30	1,379	0,879	2,962	98,60
25	400	40	86	6,009	–	26	0,860	0,792	3,369	96,62
			89	4,837	2	27	0,880	0,792	3,352	97,28
20	500	100	93	13,92	–	30	0,619	0,754	3,676	92,18
			92	6,439	2	33	0,674	0,757	3,539	96,38
10	1000	200	93	21,34	–	28	0,481	0,699	3,837	88,01
			95	13,53	2	28	0,493	0,700	3,794	92,40
1	10000	3	76	178,01	–	24	0,341	0,609	4,133	–

Результати таблиці дають змогу вибрати оптимальне співвідношення між кількістю груп, кількістю кластерів в групі та якісних оцінок результатів кластеризації та часових затрат на декомпозицію множини ключів образів.

На рис. 2 продемонстровано характер зміни функції критерію об'єднання вершин дерева (кластерів) із збільшенням рівня ієрархії дерева. На рис. 2, *а* зображено декомпозицію без групування та каскадів, на рис. 2, *б* – декомпозицію із групуванням та одним каскадом, на рис. 2, *в* – декомпозицію із групуванням та двома каскадами. Графіки дають змогу знайти оптимальне розбиття ключів за такими оцінками: 1) за кількістю кластерів; 2) за кількістю ключів в кластерах; 3) за якісними характеристиками розбиття загалом та окремими кластерами (питомі густина, об'єм, дисперсія тощо).

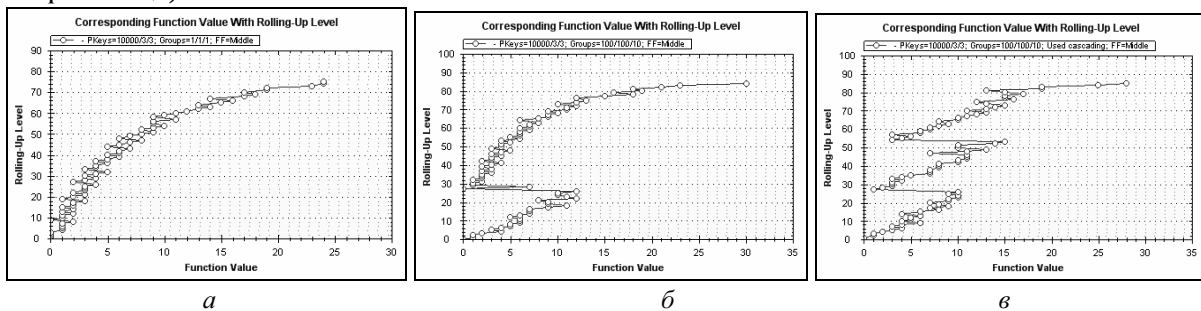


Рис. 2. Графіки функції критерію об'єднання

На рис. 3 зображено дерево згортання каскадної декомпозиції для вибірки з рівномірним законом розподілу. На ньому виділені три кластери, кожен з яких містить відповідно 3333, 3333 та 3334 листків.

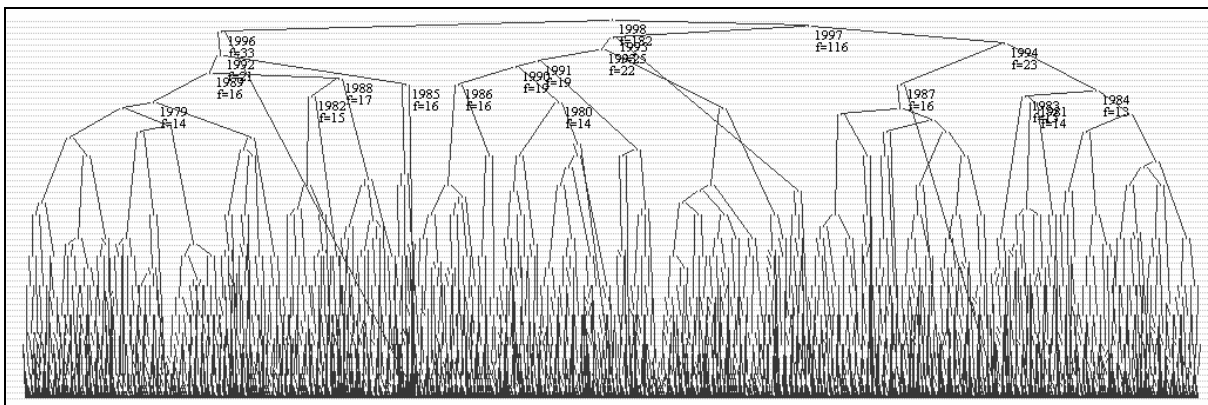


Рис. 3. Дерево згортання

Висновок. При створенні алгоритму кластеризації ключів великих вибірок для зменшення складності алгоритму необхідно розв'язати задачі точності та якості отриманих груп ключів з врахуванням практично можливих затрат часу. Запропоновано застосувати каскадну декомпозицію множин ключів (з нульового каскаду) та їх кластерів (з вищих каскадів) з метою зменшення алгоритмічної складності. Отримані результати тестування підтвердили доцільність запропонованого підходу.

1. Andy M Yip, Chris Ding, Tony F.Chan . *Dynamic Cluster Formation Using Level Set Methods*. – *IEEE Trans. on Pattern Analysis and Machine Intelligence*, vol.28, n. 6, pp.877–889, June, 2006. 2. Leo Grady, Eric L.Schwartz. *Isoperimetric Graph partitioning for Image segmentation*. – *IEEE Trans. on Pattern Analysis and Machine Intelligence*, vol.28, n. 3, pp.469–475, March, 2006. 3. M Pavan, M Pelillo. *Dominant sets and Pairwise Clustering*. – *IEEE Trans. on Pattern Analysis and Machine Intelligence*, vol.29, n. 1, pp.167–172, January, 2007. 4. Мельник Р.А., Алексєєв О.А. Кластеризація ключів образів з метою прискорення доступу до них // *Фізичне моделювання та інформаційні технології*. – 2005. – Вип. № 2 – С.172–178. 5. Мельник Р.А., Алексєєв О.А. Кластеризація ключів образів на основі декомпозиції їх множини // *Відбір і обробка інформації*. – 2006. – Вип. 24(100). – С.110–114.

УДК 519.15:621.372

О. Різник, Б. Балич, В. Парубчак
Національний університет “Львівська політехніка”
кафедра автоматизованих систем управління

ШВИДКИЙ СИНТЕЗ НЕЕКВІДИСТАНТНИХ КОНФІГУРАЦІЙ КОМБІНАТОРНИХ СТРУКТУР З ВИКОРИСТАННЯМ ЧИСЛОВИХ В'ЯЗАНОК

© Різник О., Балич Б., Парубчак В., 2007

Розглянуто швидкий синтез нееквідистантних конфігурацій комбінаторних структур на основі числових в'язанок для кодування інформації. Розроблена методика дає можливість побудови не лише одновимірних лінійок-в'язанок мінімальної довжини на послідовностях упорядкованих чисел, але й багатовимірних лінійок-в'язанок мінімальних розмірів з нееквідистантною структурою.

In the article the rapid synthesis of uneven configurations of комбінаторних structures is examined on the basis of numerical bundles. The developed method enables construction not only unidimensional lines-bundles minimum length on sequences well-organized numbers but also multidimensional lines-bundles of low-limit with a uneven structure.

Вступ. Комбінаторні конфігурації з нееквідистантною структурою, які розглядаються в роботі, є певного виду математичними моделями, за допомогою яких зручно досліджувати комбінаторні властивості багатоелементних систем будь-якої фізичної природи, наприклад, багатоелементні антени.

Ці моделі доцільно подати у вигляді множини послідовно з'єднаних між собою елементів, значення яких, як і всі можливі суми поруч розташованих елементів, різні. Подібні моделі знаходять застосування й у галузі радіоелектроніки. Фундаментальні дослідження проводяться в Південнокаліфорнійському університеті під керівництвом професора С.У. Голомба [1]. Результати досліджень були використанні в Національному геодезичному центрі Національної служби США з океану і атмосфери для поліпшення роздільної здатності віддалених радіоджерел за допомогою методу