

Изследване на клас устойчиви алгоритми за маркиране с водни знаци

Георги Върбанов, Петър Антонов

An Analysis of a Class Robust Algorithms for Watermarking: *We used some basic models and schemes imposed as a base for the development of systems for the protection of multimedia content. In the article outlining used some basic methods for embedding watermarks in widespread distribution, with a focus on sustainable (healthy) algorithms to find the right choice to reach new goals. For comparison is proposed compound algorithm containing more than one integrated into the advantages of the algorithms built into reaching a new quality and do not require the original receipt of image. I.e. this is robust blind algorithm.*

Key words: *watermark, non-blind algorithm, robust*

ВЪВЕДЕНИЕ

Последните 10-15 години се характеризират със съществено развитие и разпространение на дигиталните технологии в различни сфери на живота. Известно е, че дигиталните обекти могат да бъдат лесно модифицирани, поради което постоянно се повишават и изискванията към тяхната сигурност. Едно от най-често използваните ефективни и перспективни средства за удовлетворяване на тези изисквания е маркирането с водни знаци. Към настоящия момент в това направление са разработени множество алгоритми, които могат да бъдат класифицирани по различни признаци. Например, според необходимостта от наличие на оригиналното изображение при извличане на водния знак, алгоритмите за маркиране се разделят на [1]: *слепи (blind) алгоритми*, при които за откриване на водния знак не е необходимо да се притежава оригиналното изображение, *полуслепи (semi-blind)* и *неслепи (non-blind)* - където за откриване на водния знак е необходимо оригиналното изображение. Съществуват класификации и по други класификационни признаци, като модифициращ метод, тип на средата за маркиране и др. Независимо от множеството разработки обаче, задачата за усъвършенстване на алгоритмите за маркиране с водни знаци и за създаване на нови такива продължава да е особено актуална. В тази връзка, в настоящия доклад се прави анализ на известни устойчиви алгоритми и на тяхна база се представя усъвършенстван слеп (blind) алгоритъм с подобрени качества.

ОПИСАНИЕ НА АЛГОРИТЪМА

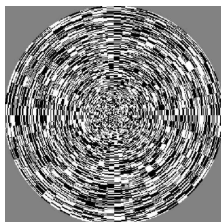
Представеният алгоритъм комбинира елементи на известни устойчиви слепи и неслепи алгоритми от [3,4,8 и др.]. От [8] е използвана идеята за генериране на себеподобен воден знак, Себеподобието в случая осигурява изключително висока устойчивост срещу много видове атаки, като изкривявания, линейно и нелинейно филтриране, ротация, изрязване, скалиране, JPEG компресия. Използва се Wavelet преобразуване за да се получи висока надеждност и визуална нерушимост на изображението от направените промени. Извличането на водния знак става чрез статистическо изчисление на нормализирана зависимост на маркираното изображение от водния знак и сравнение с предврително избрана стойност, такава че да не се допуска фалшива тревога за наличие на воден знак, или да не се открива вграден воден знак. Основната част от настоящия алгоритъм е генерирането на водния знак, който е така подбран, че самата му структура да осигурява изключителна защита срещу скалиране и изрязване на маркираното изображение. Водният знак представлява $W_M(r, \theta)$ -кръгово-симетричен сигнал. Този сигнал се описва с функция от (r, θ) , която има следните крайни стойности:

$$W_M(r, \theta) = \begin{cases} 0, & \text{ако } (r < r_{\min} \text{ или } r < r_{\max}) / \text{или ако } r_{\min} < r < r_{\max} \\ \pm \alpha & \end{cases} \quad (1)$$

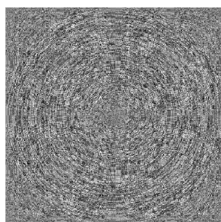
където

$r = \sqrt{n_1^2 + n_2^2}$ (n_1 и n_2 - са пространствените координати на пиксела); $r_{\min}$ и $r_{\max}$ са минималният и максималният радиус, които описват големината на водния знак; $\theta = -\arctan\left(\frac{n_2}{n_1}\right)$; α е цяло число, представляващо нивото на вграждане.

За да бъде функцията $W_M(r, \theta)$ устойчива на компресия или филтриране, сигналът се разделя допълнително на s сектора, всеки с дъга от $s/360^\circ$. Всички пиксели в един сектор, имащи един и същ радиус, се приравняват на $-\alpha$ или $+\alpha$, в зависимост от псевдослучаен генератор на случайни числа, инициализиран със начална стойност k . Сигналът, описващ самоподобния воден знак има вида, показан на фиг. 1



фиг.1



фиг.2

В случая, това разделяне на сектори се прилага за да се осигури противодействие срещу атаките с ротация. Откриването на водния знак при ротация на ъгли, по-малки от $s/360^\circ$, се извършва без да се завърта водния знак. Когато изображението е завъртяно на повече от $s/360^\circ$, тогава за откриването на водния знак е необходимо неговото завъртане на ъгъл, кратен на $s/360^\circ$. На следващата стъпка е необходимо да се вгради себеподобие в отношение с декартовата равнина. Сигналът $W_M(n_1, n_2)$ се скалира (с използване на интерполация със съседни стойности), след което търсенето се измества 4 пъти спрямо центъра $r=0$. Това позволява да се открие видния знак при преоразмеряване на картината от 1/8 до 8 пъти. Така крайният воден знак се състои от 4 слоя :

$$W(n_1, n_2) = \sum_{f=0}^3 \sum_{i=0}^{2^f-1} \sum_s W_m \left(2^f n_1 + i \left[\frac{r_{\max}}{2^f} \right], 2^f n_2 + i \left[\frac{r_{\max}}{2^f} \right] \right) \quad (2)$$

Вижда се, че водният знак е вместен в квадрат с размери $(2r_{\max} \times 2r_{\max})$. Параметрите, които са необходими за детекцията на водния знак са $[k, \alpha, r_{\min}, r_{\max}, s]$

Крайният резултат за генерирания воден знак е показан на фиг. 2. Следващата необходима стъпка е вграждане на този знак.

Вграждането се извършва с адитивен алгоритъм, разгледан подробно в [3].и [4] В най-общи линии този алгоритъм може да бъде описан по следния начин:

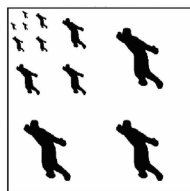
Изображението се декомпозира 2 пъти с използване на Haar Wavelet трансформация. Така от двете нива на Haar преобразуването се получават общо 8 нива на декомпозиция, които се обозначават с LL1, LH1, HL1, HH1, LL2, LH2, HL2, HH2 (LL1 - нискочестотни коефициенти от първо ниво, LH1 и HL1 - средночестотни коефициенти от първо ниво, HH1 - високочестотни коефициенти от първо ниво, LL2

– нискочестотни коефициенти от второ ниво и т.н. Тъй като най-много информация за картината се съхранява в нискочестотните коефициенти, за да може да се постави водния знак така, че да не е видим, коефициентите на знака се наслагват към нискочестотните и средночестотните коефициенти. Освен това се прави допълнителна маскировка на вмъкнатите битове, така че да не са видими и да не могат да се открият чрез сканиране на честотите на изображението или по някакъв друг метод за анализ на цифрови сигнали. Както вече бе пояснено, първата стъпка към това е генерирането на водния знак, който се получава от налагането на функцията $WM(x, y)$, която има крайни стойности $(-1, 1)$ и средна значеща стойност 0. Тази функция представлява Гаусов шум, който се налага към изображението и прави знака практически неоткриваем, тъй като генерирането му става с помощта на псевдослучаен генератор на случайни числа. Допълнителното маскиране, което се прилага към картината, е резултат от прилагането на алгоритъма ОПИ (Области, които Представляват Интерес). Това са тези области, които съдържат по-голямата част от информацията. Те се определят с помощта на принципа за MTWC - многопрагово wavelet кодиране (Multi-threshold Wavelet coding). Този принцип описва извличането от най-значещите коефициенти на даден вектор, като най-значещ означава такъв коефициент, който има най-голяма амплитуда. Маркират се по 2 съседни коефициента, като се добавя/изважда стойността на водния знак, реципрочно от по-големия и по-малкия коефициенти, с цел намаляване на вложения шум в изображението. Общо казано тези коефициенти не се променят значително под въздействието на атаки от типа на компресия или геометрическа обработка. Ако тези коефициенти бъдат променени значително под някакво външно въздействие, то тогава изображението би било много по-различно от това, в което сме вградили водния знак.

Определянето на най-значещите области става чрез селекция на най-значещите битове. Реално се избира само тази част от картината, която съдържа най-важната информация, която в примера от фиг. 3 е фигурата на сърфиста.



фиг. 3



фиг. 4

След прилагане на Wavelet преобразуването, изображението, в което ще се влагат стойностите на водните знаци ще съдържа само стойностите на областта, в която се намира сърфиста.

Алгоритъмът за определяне на коефициентите се изпълнява в следните стъпки:

- инициализация: установяване на праговата стойност T_s за всяка честота, се равнява на половината от максималната абсолютна стойност по всички коефициенти от тази област;
- избиране на честотна област, в която максималният коефициент има стойност $p \cdot T_s$, където p е мащабиращ коефициент, указващ силата на водния знак. За тази област се селектират всички коефициенти, които имат стойности по-големи от избраната честота T_s , и се маркират като значещи;

- водният знак се вмъква в избраните коефициенти от предната стъпка и се установява нова стойност за $T_{s\text{new}} = T_s / 2$.

- повтарят се стъпките от 2 до 4, докато се вмъкнат всички коефициенти от водния знак.

Формулата, която описва начина на вмъкване на водния знак има следният вид:

$$C_s^*(x, y) = C_s(x, y) + Y \cdot \beta_s \cdot T_s \cdot W_k, \quad (3)$$

където C_s е избраният оригинален коефициент; C_s^* е маркираният коефициент; T_s е праг на вмъкване за текущата честота; W_k е стойност на водния знак в диапазона $\{-\alpha, \alpha\}$; Y и β_s са коефициенти на квантуване.

Стойността на Y се избира адаптивно, така че да не се промени видимо картината. Този коефициент позволява да се увеличава/намалява маскирането на водния знак и да се намалява/увеличава сигурността на маркирането. Той се избира в интервала (0, 1). Параметърът β_s се използва за контрол на избора на коефициентите в дадената честота - колкото по-голяма е стойността на $C_s \times \beta_s$, толкова е по-голяма честотата, в която могат да се вмъкват коефициентите на водните знаци W_i . За извличане на водния знак се използва следната формула:

$$E_{s,k}^*(x, y) = C_{s,k}^*(x, y) - C_{s,k}(x, y), \quad (4)$$

където: C_s^* е коефициента, който е извлечен от вероятно атакуваното изображение I^* ; C_s е коефициента от оригиналното изображение I .

Детекцията на водния знак се извършва чрез пресмятане на подобие между C_s^* и C_s) по формулата

$$SIM(I^*, I) = \frac{\sum_{k=1}^{N_w} E_{s,k}^*(x, y) \cdot E_{s,k}(x, y)}{\|E_{s,k}^*(x, y)\| \|E_{s,k}(x, y)\|}, \quad (5)$$

където N_w е общия брой на символите от водния знак; $E_{s,k}(x, y)$ е оригиналният воден знак и $E_{s,k}^*(x, y)$ е водния знак, извлечен от коефициентите на атакуваното изображение.

Важно е да се отбележи, че се използва скалярно произведение на два вектора, нормализирано с техните нормални стойности, като мярка за подобие, което се измерва с косинусовата функция на ъгъла между двата вектора. Всяко подобие, по-малко от 0, се третира като нула, тъй като тогава тези два вектора са с противоположни посоки. Максималната стойност на това подобие е 1, когато имаме неатакувано изображение и водният знак е извлечен изцяло.

Въпреки, че водните знаци, генерирани от псевдослучайните последователности не могат да дадат по-добра корелация от случайните числа с равномерно разпространение, енергията на различните водни знаци може да се направи подобна, така че да се разпръсне в някакви граници. При това, водните знаци, генерирани с различни инициализиращи стойности на псевдослучайния генератор могат да имат една и съща степен на устойчивост срещу атаки.

Описаното дотук представя алгоритъм за non-blind водно маркиране, което изисква наличието на оригиналното изображение за извличане на водния знак. Тъй като избраният за реализация алгоритъм е blind, то се прилагат някои трансформации на алгоритъма, за да може да се прави детекция, без наличие на оригиналното изображение. Най-трудният проблем при подобни схеми е да се определят коефициентите, в които е вмъкнат водния знак, както и да се отделят коефициентите на самия воден знак. Основната идея на blind алгоритмите е да се изрежат получените коефициенти на базата на генерирани псевдооригинални

стойности, след което водният знак да се наложи върху тези коефициенти за да се получи защитено изображение.

Нека $C_s(x, y)$ е избраният коефициент от честота s с текуща прагова стойност T_s , така че $T_s \leq C_s(x, y) < 2 \times T_s$. Тогава защитената стойност на C_s ще бъде:

$$C_{s,k}^*(x, y) = \text{sign} \times \Delta_p(C_s(x, y)) + \alpha \beta_s T_s W(x, y), \quad (6)$$

където $W(x, y)$ е стойността на водния знак; sign е знак на стойността на $C_s(x, y)$ и операторът Δ е дефиниран като:

$$\Delta_p(C_s(x, y)) = (1 + 2p\alpha)T_s, \quad (7)$$

където p е цяло число между 1 и $(2\alpha)^{-1}$.

Тогава разстоянието $DIS_{s,p}(x, y)$ между $\Delta_p(C_s(x, y))$ и $C_s(x, y)$ се определя от следната формула:

$$DIS_{s,p}(x, y) = |\Delta_p(C_s(x, y)) - C_s(x, y)|, \quad (8)$$

където p се определя като $p = \arg \min_p DIS_{s,p}(x, y)$

След намирането на p , можем да запишем

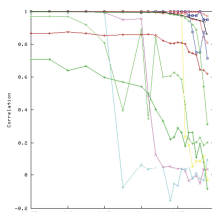
$$DIS_{s,p}(x, y) \leq \alpha \times T_s \quad (9)$$

Детекцията на водния знак се изчислява по дадената по-горе формула (5), където $E_{s,k}(x, y)$ се замества с

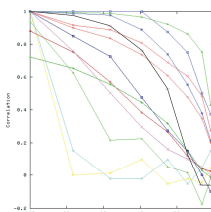
$$E_{s,k}^*(x, y) = C_{s,k}^*(x, y) - \text{sign} \times A_p C_{s,k}^*(x, Y), \quad (9)$$

където Δ_p е равно на $(1 + 2p\alpha)T_s$, а T_s се изчислява от C^* .

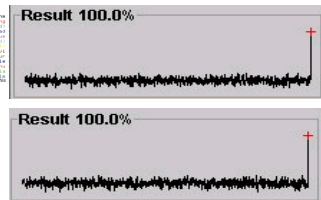
Тъй като T_s^* се определя на базата на най-големия коефициент в дадена честота s и в него не се вмъква никаква допълнителна информация, то може да се счита, че той ще бъде максимално близък до оригиналния. Ако се вземе за даден коефициент $\alpha \cdot p \leq 1.0$ и ако стойността на коефициента $C_{s,k}^*(x, y)$ е била променена с по-малко от $\alpha \times T_s$, тогава водният знак в $C_{s,k}^*(x, y)$ може да бъде извлечен изцяло. По-големи стойности на α могат да осигурят по-голяма устойчивост на водния знак, но също така и по-голяма стойност на шума, генериран от водния знак $PSNR$. В допълнение към защитата чрез маркиране с воден знак, методът с маркиране в ОПИ може да се използва за маркиране на данни в големи бази от данни с цел полесното извличане на информацията. Например, при архивирането на изображения най-важното е да се реализира някакъв много бърз алгоритъм за управление на данните, тъй като в този случай имаме изключително голямо количество информация. Чрез използване на традиционното индексване на бази данни е изключително трудно да се индексират данните според тяхното местоположение, големината и връзки между обектите.



фиг.5



фиг.6



фиг.7

На фиг. 5 и 6 са представени част от експерименталните резултати за сравнение между известни blind и non-blind алгоритми и представеният в доклада комбиниран алгоритъм, маркиран с (this). Приведената на фиг. 7 част от изследванията на комбинирания алгоритъм показва неговата абсолютна устойчивост срещу медианна (3x3) и гаусова (3x3) атаки. Всички експериментални резултати са получени при атаки на маркираните изображения със StirMark3.

ЗАКЛЮЧЕНИЕ

От анализа на различните стандартни blind и non-blind алгоритми се потвърждава известния факт, че non-blind алгоритмите имат значително по-добри резултати от blind алгоритмите. Освен това, резултатите показват, че предложеният в доклада комбиниран алгоритъм се характеризира с много-близки до non-blind алгоритмите параметри, като в същото време отпада необходимостта от запазване в база данни на оригиналното изображение.

ЛИТЕРАТУРА

- [1] Kundur, D., D. Hatzinakos. A Robust Digital Image Watermarking Method using Wavelet-based Fusion. //Proc. of IEEE Intern. Conf. on Image Processing (ICIP'97). Santa Barbara, California, 1997, pp. 544-547 [<http://www.ece.tamu.edu/~deepa/pub.html>].
- [2] Hartung, F., J. K. Su, B. Girod. Spread Spectrum Watermarking: Malicious Attacks and Counterattacks in Security and Watermarking of Multimedia Contents. //Proc. of SPIE 3657, 1999, pp. 147-158.
- [3] Xiang-Gen, X., C. G. Boncelet, G. R. Arce. A Multiresolution Watermark for Digital Images. //Proc. of IEEE International Conference on Image Processing (ICIP '97), Vol. 1, Santa Barbara, California, 1997 [<http://www.cpsy.sbg.ac.at/~pmeerw/watermarking/>].
- [4] Xiang-Gen, X., C. G. Boncelet, G. R. Arce. Wavelet Transform Based Watermark for Digital Images. //Optics Express, 3, December 1998.
- [5] Petitcolas, F., R. Anderson, M. Kuhn. Attacks on Copyright Marking systems. //Second Workshop on Information Hiding, 1998.
- [6] Lu, C., S. Huang, C. Sze, H. Liao. Cocktail Watermarking for Digital Image Protection. IIS, Taiwan, [<http://citeseer.ist.psu.edu/article/lu00cocktail.html>].
- [7] Fabien, A., P. Petitcolas. Watermarking schemes evaluation. //IEEE Signal Processing, 2000, Vol. 17, No 5, pp. 58-64.
- [8] Tsekeridou, S., I. Pitas, "Embedding Self-Similar Watermarks in the Wavelet Domain" 2000 IEEE Int. Conf. on Acoustics, Systems and Signal Processing (ICASSP'00), vol. IV, pp. 1967-1970, Istanbul, Turkey, 5-9 June 2000

За контакти:

Доц. д-р инж. Петър Цветанов Антонов, Катедра "Компютърни науки и технологии", ТУ - Варна, тел: 052-383 320, e-mail: peter.antonov@ieee.org

Гл.ас. инж. Георги Върбанов, Катедра "Компютърни науки и технологии", ТУ - Варна, тел: 052-383 614, e-mail: gvarbanov@gmail.com

Докладът е рецензиран.