

UDC 658

<https://doi.org/10.32362/2500-316X-2021-9-5-7-13>

RESEARCH ARTICLE

The structure of the local detector of the reprint model of the object in the image

Alexander A. Kulikov®*MIREA – Russian Technological University, Moscow, 119454 Russia*® Corresponding author, e-mail: tibult41@gmail.com

Abstract. Currently, methods for recognizing objects in images work poorly and use methods that are intellectually unsatisfactory. The existing identification systems and methods do not completely solve the problem of identification, namely, identification in difficult conditions: interference, lighting, various changes on the face, etc. To solve these problems, a local detector for the reprint model of the object in the image is developed and described. A transforming autocoder (TA), a model of a neural network, has been developed for a local detector. This neural network model is a subspecies of the general class of neural networks of reduced dimension. The local detector is able, in addition to determining the modified object, to determine the original shape of the object as well. A special feature of TA is the representation of image sections in a compact form and the evaluation of the parameters of the affine transformation. The transforming autocoder is a heterogeneous network (HS) consisting of a set of networks of smaller dimension, called a capsule. Artificial neural networks should use local capsules that perform some rather complex internal calculations on their inputs, and then encapsulate the results of these calculations in a small vector of highly informative outputs. Each capsule learns to recognize an implicitly defined visual object in a limited area of viewing conditions and deformations, and it outputs both the probability that the object is present in its limited area, and a set of «instance parameters» that can include the exact pose, lighting, and deformation of the visual object relative to an implicitly defined canonical version of this object. The main advantage of capsules that output instance parameters is a simple way to recognize entire objects by recognizing their parts. The capsule can learn to display the pose of its visual object in a vector that is linearly related to the «natural» representations of the pose that are used in computer graphics. There is a simple and highly selective test for whether visual objects represented by two active capsules A and B have the correct spatial relationships for activating a higher-level capsule C. The transforming auto-encoder solves the problem of identifying facial images in conditions of interference (noise), changes in illumination and angle.

Keywords: neural network, image recognition, pattern recognition, identification model

• Submitted: 25.03.2021 • Revised: 31.03.2021 • Accepted: 26.05.2021

For citation: Kulikov A.A. The structure of the local detector of the reprint model of the object in the image. *Russ. Technol. J.* 2021;9(5):7–13. <https://doi.org/10.32362/2500-316X-2021-9-5-7-13>

Financial disclosure: The author has no a financial or property interest in any material or method mentioned.

The author declares no conflicts of interest.

НАУЧНАЯ СТАТЬЯ

Структура локального детектора модели репринта объекта на изображении

А.А. Куликов [®]

МИРЭА – Российский технологический университет, Москва, 119454 Россия

[®] Автор для переписки, e-mail: tibult41@gmail.com

Резюме. Задача распознавания объектов на изображениях является актуальной в настоящее время, поскольку существующие системы и методы не решают полностью проблему идентификации в сложных условиях: помехи, освещение, различные изменения на лице и т.д. С целью решения этой задачи разработан и описан локальный детектор для модели репринта объекта на изображении. Для локального детектора разработан трансформирующий автокодер (ТА) – модель нейронной сети. Данная модель является подвидом общего класса нейронных сетей снижения размерности. Локальный детектор способен, помимо определения измененного объекта, также определить и изначальную форму объекта. Особенностью ТА является представление участков изображения в компактном виде и проведение оценки параметров аффинной трансформации. Трансформирующий автокодер представляет собой гетерогенную сеть (ГС), состоящую из множества сетей меньшей размерности, называемых капсулами. Искусственные нейронные сети должны использовать локальные капсулы, которые выполняют некоторые довольно сложные внутренние вычисления на своих входах, а затем инкапсулируют результаты этих вычислений в небольшой вектор высокоинформативных выходов. Каждая капсула учится распознавать неявно определенный визуальный объект в ограниченной области условий просмотра и деформаций. Она выводит как вероятность того, что объект присутствует в своей ограниченной области, так и набор «параметров экземпляра», которые могут включать точную позу, освещение и деформацию визуального объекта относительно неявно определенной канонической версии этого объекта. Главное преимущество капсул, выводящих параметры экземпляра, заключается в простом способе распознавания целых объектов путем распознавания их частей. Капсула может научиться выводить позу своего визуального объекта в вектор, линейно связанный с «естественными» представлениями позы, которые используются в компьютерной графике. Существует простой и высокоселективный тест на то, имеют ли визуальные объекты, представленные двумя активными капсулами, правильные пространственные отношения для активации капсулы более высокого уровня. Трансформирующий автокодер решает проблему идентификации лиц на изображениях в условиях помех (шумности), изменения освещенности и ракурса.

Ключевые слова: нейронная сеть, распознавание изображений, распознавание образов, модель идентификации

• Поступила: 25.03.2021 • Доработана: 31.03.2021 • Принята к опубликованию: 26.05.2021

Для цитирования: Куликов А.А. Структура локального детектора модели репринта объекта на изображении. *Russ. Technol. J.* 2021;9(5):7–13. <https://doi.org/10.32362/2500-316X-2021-9-5-7-13>

Прозрачность финансовой деятельности: Автор не имеет финансовой заинтересованности в представленных материалах или методах.

Автор заявляет об отсутствии конфликта интересов.

INTRODUCTION

Existing systems and methods for object recognition in images do not completely solve the problem of identification, specifically the identification under complex conditions: interference, illumination, changes in faces, changes in wide-angle shooting, etc. Presently, employing sets of nonlinear

functions for activation of used neurons represents a laborious and inaccurate process. This statement—supported by a large number of publications [1–13] devoted to this problem—says that the problem is topical and is not solved yet. Corresponding methods, algorithms, and systems either require high computational power, or employing a programmed, continuously operating memory device in “clever”

cameras that leads to growth in expenses of the entire system.

LOCAL DETECTOR

To solve the above problem of identification of facial images (and not only those), a local detector (LD) is developed for the model of the object reprint in the image [14]. LD is an elementary unit of the model of the object reprint (MOR) in the image. For the local detector, a transforming autocoder (TA)—the model of a neural network—has been developed. This model represents a subtype of a general class of neural networks of reduced dimension. Besides identification of a changed object, LD is capable of determining initial shape of the object. One of the features of TA is its ability represent parts of the image in a compact form and evaluate the parameters of affine transformation.

The transforming autocoder represents a heterogeneous network (HN) that consists of a set of lower dimensions—capsules.

Definition/Features of a capsule:

- All capsules of the transforming autocoder have the same structure.
- Each capsule encapsulates the visualization of the object's image.

TA is a neural network with the ability to learn, for which a “method of reverse propagation of the error” is directly employed. Input values of the autocoder are used for normalization. For the neural network under consideration, the function has a simple form $c = f(x) = x$. When using a transforming autocoder, it is additionally necessary to apply the restriction of “bottle neck” in one of layers with lower number of neurons than that in the input layer.

Thus, the neurons within such a type of the layer represent a reprint of data. As distinguished from the main components' method, using a set of layers of a transforming auto-coder and nonlinear functions of activation of neurons is compact and accurate.

Here is an example. When a set of data (an image in our case) is delivered to an input and is represented as a set of small images having the size $x \in R^{28 \times 28 = 784}$, then their reprint can be represented as a hidden layer having the size of the order of 30, that is $c = f(x) = R^{30}$. In each capsule, there is one crucial neuron with values (0, 1), that corresponds to the fact that the object is in the image.

Some systems of computer imaging make use of histograms of oriented gradients as “visual words” and simulate spatial distributions of those elements with the help of a rough spatial pyramid. Such methods can correctly recognize objects without correct knowledge of their location—the ability that is used to diagnose the human's brain damage. Artificial neural networks make use of schemes of weight distribution with manual

encoding to reduce the number of free parameters and reach a local translational invariant by subsampling the activation of local pools of translated replicas of the same core. After several stages of subsampling in a converging network, high-level objects have greater uncertainty in their poses.

Artificial neural networks should employ local capsules, which perform some rather complex internal calculations at their inputs, and then encapsulate the results of those calculations in a small vector of highly informative outputs. Each capsule learns to recognize implicitly the determined visual object within a limited region of views and deformations. The capsule's outputs are the probabilities for the object to occupy its limited region and a set of “parameters of a copy” that can include accurate pose, illumination and deformation of a visual object relatively implicitly determined canonical version of this object. When a capsule operates as needed, the probability of presence of a visual entity represents a local invariant; it does not change when the entity moves over the set of possible realizations in encapsulated limited region. The parameters of the copy are “equivariantly”: as viewing conditions and motions of the object change over the external set, the parameters of the copy change by the same value since they represent internal coordinates of the object in the external set [1].

One of major advantages of capsules delivering explicit parameters of a copy consists in a simple way of recognizing entire objects by recognizing their parts. If a capsule can learn to deliver a pose of its virtual object to a vector, that is linearly coupled with an used in computer graphics “natural” representation of a pose, there is a simple and highly selective test for checking of whether visual objects (represented by two active capsules A and B) have accurate spatial relationship for activating a capsule of higher level C. Let us suppose that output data of capsule A represented by matrix T_A that yields the transformation of coordinates between a canonical visual entity and the actual copy of that entity determined by capsule A. If T_A is multiplied by a functional of coordinates transformation “a part-to-a whole” T_{AC} , which couples canonical visual entities A and C, then a predicted matrix T_C can be obtained. In a similar way, we can use T_B and T_{BC} for obtaining another prognosis. If these predictions agree well, then the copies determined by capsules A and B are in correct spatial relationship enabling to activate capsule C; the average value of the prediction informs us of how the greater visual entity—represented by C—transforms relatively canonical visual entity C. For example, if A is a mouth and B is a nose, then each of them can predict the position of the face. If these predictions coincide, then both the mouth and the nose must be in a correct spatial relation to form a face. An interesting feature of

this approach for recognizing a form is that the relation “a whole-to a part” is invariant and is represented by weighting matrices, whereas the ensemble of parameters of a copy of objects and their parts being observed at the given moment is equivariant and is represented by neural activities.

To obtain “a whole-to a part” hierarchy, the capsules, which realize parts of the lowest level in this hierarchy, must extract from intensities of pixels explicit parameters of the pose.

After intensities of pixels have been transformed to output data of the given set of active capsules of the first level, each of them yields explicit representation of the pose of its visual object, visual objects can be recognized with the help of active capsules of lower level.

Let us consider a neural network of direct communication shown in Fig. 1.

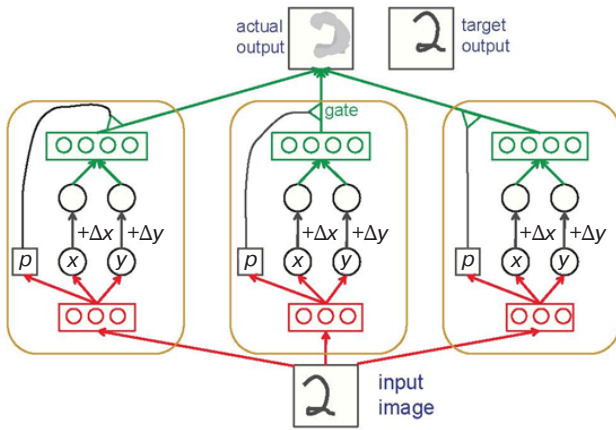


Fig. 1. Three capsules of a transforming autocoder that simulates translations

Each capsule in Fig. 1 has three recognizing blocks and four generating blocks. The weights at connections are learned through reverse propagation of a discrepancy between actual and targeted outputs. The

network is determined, and as soon as it is learned it can receive the image as input data and desired shifts Δx and Δy and generate the shifted output. Each capsule has its own logical “recognizing blocks,” which act as a hidden layer for the calculation of outputs, which are three numbers, x , y , and p . The capsule will send these outputs to higher levels of the visual system (p is the probability that the visual entity of the capsule is presented in the input image). In its turn, the capsule has its own generators; values $x + \Delta x$ and $y + \Delta y$ being input data for these blocks (where x and y are input and output data for one capsule). The capsules research generative units with projected fields of strong localization (Fig. 2). We describe the state for each of them by the following activation functions:

$$\begin{aligned} H_r &= \sigma(\mathbf{W}_{xh}x + b_r) \in (0,1)^{N_r}, \\ c &= \mathbf{W}_{hc}H_r + b_c \in R^2, \\ c' &= c + s \in R^2, \\ p &= \sigma(\mathbf{W}_{hp}H_r + b_p) \in (0,1), \\ H_g &= \sigma(\mathbf{W}_{c'g}c' + b_g) \in (0,1)^{N_g}, \\ y &= p\mathbf{W}_{hy}H_g \in R^{784}. \end{aligned} \quad (1)$$

Note that if every capsule receives nine real inputs, which are processed as a matrix with 3×3 dimension, then TA can be learned to predict complete 2D affine transformation (translation, rotation, scaling, and shift). Transformation matrix $\mathbf{T}$ is applied to the output of capsule A to obtain matrix $\mathbf{T}_A$. Then the elements of $\mathbf{T}_A$ are used as input data for generation blocks when $\mathbf{T}_A$ for the prognosis of targeted output image.

A criterion of rarefaction is supposed to be used as an auxiliary condition. Employing this criterion is effective for TA. Note, that the condition of rarefaction

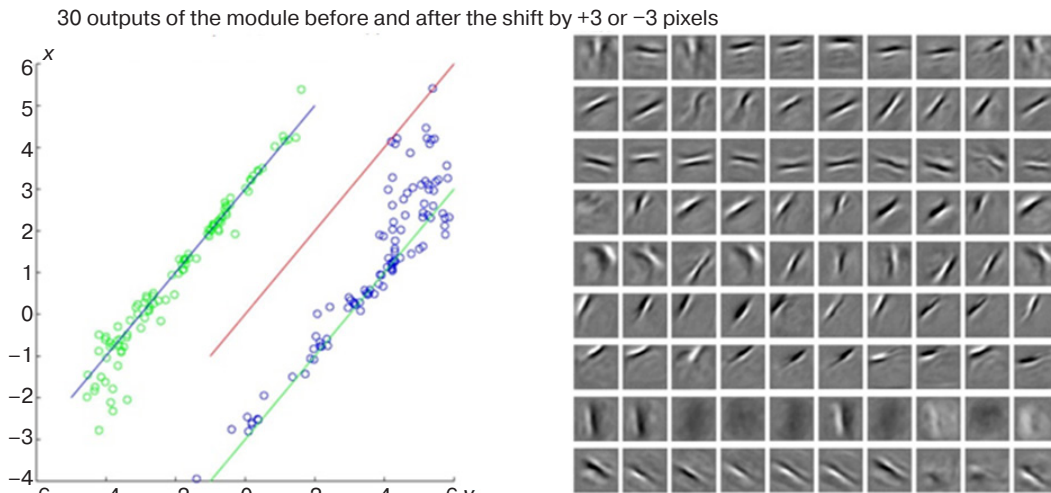


Fig. 2. Shifts in values of x and y

must impose restriction on neurons of a generating and recognizing layer. Therefore, for employing TA as a constituent of the model we should extend a formulation for the transformation beyond the limits of two-dimensional translation. For the rarefaction condition, the information divergence of the facial image identification model already developed by the author is applied. The Kullback–Leibler divergence formula is as follows:

$$S = \sum_{j=1}^{L_n} KL(s \| \hat{s}_j) = \sum_{j=1}^{L_n} s \log \frac{s}{\hat{s}_j} + (1-s) \log \frac{1-s}{1-\hat{s}_j}, \quad (2)$$

where L_n is the number of neurons, $\hat{s}_j = \frac{1}{m} \sum_{i=1}^m (a_j^{(Ln)} x_j)$ is the average value of activation, and s is the rarefaction parameter.

By setting the value of parameter s to be proportionally small, we can limit the average activation of a neuron. Having obtained independent signs, we can also change the impact of this parameter on the operation of the heterogeneous network.

For further optimization of the parameters, it is necessary to employ a price function TA with weights D and v , which can be written as follows:

$$J_s(\mathbf{D}, \mathbf{v}) = J(\mathbf{D}, \mathbf{v}) + \beta \sum_{i=1}^L \sum_{j=1}^{L_i} KL(s \| \hat{s}_j), \quad (3)$$

where β is a meta-parameter, $\mathbf{D}$ and $\mathbf{v}$ are general matrices of weights.

An additional parameter should be introduced into the algorithm of the error backpropagation. The error determined by the backpropagation method represents an expression for some layer of neuron network l .

$$\delta_i^{(l)} = \left(\left(\sum_{j=1}^L \mathbf{w}_{ji}^{(l)} \delta_j^{(l+1)} \right) + \beta \left(-\frac{s}{s_j} + \frac{1-s}{1-s_j} \right) \right) f'(z_i^l), \quad (4)$$

where z_i^l is the argument of a function for activation of the i th neurons in layer l .

This parameter represents the rarefaction criterion. Values s_j depends on D and v as the average activation of neuron j .

As compared to the methods and algorithms investigated in [1–3, 5, 6, 9–13, 15], TA (neural network) has demonstrated the best result in identification of facial images under different shooting conditions. Table 1 shows the results of identification at different angles of shooting.

Table 1. The results of identification for different angles

Angle	POSIT, %	SVM, %	Author's (TA), %
(0°, 15°)	82 ± 4	85 ± 2	99 ± 4
(15°, 30°)	80 ± 3	81 ± 3	98 ± 3
(30°, 45°)	79 ± 4	80 ± 4	97 ± 3
(45°, 60°)	81 ± 5	82 ± 4	98 ± 4

Table 2 shows the results of identification for different levels of illumination.

Table 2. The results of identification for different levels of illumination

Illumination, %	POSIT, %	SVM, %	Author's (TA), %
25	35 ± 2	15 ± 2	88 ± 2
50	61 ± 5	47 ± 2	98 ± 2
75	70 ± 2	68 ± 4	98 ± 1
100	99 ± 1	99 ± 1	99 ± 1

Table 3 shows the results of identification under different levels of interferences and noise in the image. We define noise as losing of the image sharpness with magnification. As for interferences, these are different interferences when producing the image as well as the appearance of additional features in the facial image: glasses, moustache, makeup, etc.

Table 3. The results of identification under different interferences and noise in the image

Parameters	POSIT, %	SVM, %	Author's (TA), %
Noise	84 ± 2	92 ± 2	97 ± 2
Interference	89 ± 5	83 ± 2	99 ± 1

The presented results demonstrate that TA employing a local detector is less responsive to a change in the position of a facial image, illumination, and interference (noise).

CONCLUSIONS

The proposed method can be extended for identification of three-dimensional objects and is promising for analyzing a combination of local spatial structures and hypothetical 3D model of the object. The described LD (the element of the MOR model) solves the problem of the stability of identification of facial images under conditions of interference (noise), change in illumination, and different angles.

REFERENCES

1. Parfinovich S.N. Algorithms of face recognition for identity verification by image. In: *"Molodoi issledovatel': vyzovy i perspektivy": sb. mat. CXIV Mezhdunarodnoi nauchno-prakticheskoi konferentsii* (Proceedings CXIV International Scientific and Practical Conference "Young Researcher: Challenges and Prospects"). Moscow: Internauka; 2019, p. 115–163. (in Russ.).
2. Akhmedov A.A., Sagidov G.S., Kurbanismaïlov G.M. Algorithm of face recognition based on the Viola–Jones method. In: *"Molodoi issledovatel': vyzovy i perspektivy": sb. mat. CXVIII Mezhdunarodnoi nauchno-prakticheskoi konferentsii* (Proceedings CXVIII International Scientific and Practical Conference "Young Researcher: Challenges and Prospects"). Moscow: Internauka; 2019, p. 270–274. (in Russ.).
3. Pentland A., Choudhary T. Face recognition for smart environments. *Otkrytye sistemy = Open Systems Publications*. 2000;03 (in Russ.). Available from URL: <https://www.osp.ru/os/2000/03/177939>
4. Gorelik A.L., Gurevich I.B., Skripkin V.A. *Sovremennoe sostoyanie problemy raspoznavaniya: Nekotorye aspekty (The current state of the recognition problem: Some aspects)*. Moscow: Radio i svyaz'; 1985. 161 p. (in Russ.).
5. Samal D.I., Frolov I.I. Algorithm of preparation of the training sample using 3D face modeling. *Sistemnyi analiz i prikladnaya informatika = System analysis and applied Information science*. 2016;4:17–23 (in Russ.). Available from URL: <https://sapi.bntu.by/jour/article/view/128/105>
6. Zavalov R.A., Garaev R.A. Implementation of the Viola–Jones algorithm on a microcontroller with limited resources. *Nauka i obrazovanie segodnya = Science and Education Today*. 2018;6(29):18–23 (in Russ.). Available from URL: <https://cyberleninka.ru/article/n/realizatsiya-algoritma-violy-dzhonsa-na-mikrokontrollere-s-ogranichennymi-resursami/viewer>
7. Baldin A.V., Eliseev D.V. Multidimensional matrix algebra for adapted data model processing. *Nauka i obrazovanie: nauchnoe izdanie MGTU im. N.E. Bauman = Science and Education of Bauman MSTU*. 2011;7:4 (in Russ.). Available from URL: <http://technomag.edu.ru/doc/199561.html>
8. Korotkov A. Database index for approximate string matching. In: *Proceedings of the 4th Spring/Summer Young Researchers' Colloquium on Software Engineering. SYRCoSE '10*. 2010, p. 136–140. <https://doi.org/10.15514/syrco-se-2010-4-27>
9. Kononykhin I.A., Ezhov F.V., Martynyuk R.A., et al. Implementation of a face recognition and tracking system. *Molodoi uchenyi = Young Scientist*. 2020;28(318):8–12 (in Russ.). Available from URL: <https://moluch.ru/archive/318/72492/>
10. Hinton G.E., Krizhevsky A., Wang S.D. Transforming auto-encoders. In: Honkela T., Duch W., Girolami M., Kaski S. (Eds.). *Artificial Neural Networks and Machine Learning – ICANN 2011. ICANN 2011. Lecture Notes in Computer Science*. Springer, Berlin, Heidelberg; 2011. V. 6791. P. 44–51. https://doi.org/10.1007/978-3-642-21735-7_6
11. Alghaili M., Li Z., Ali H.A.R. FaceFilter: Face identification with deep learning and filter algorithm. *Scientific Programming*. 2020;1–9. <https://doi.org/10.1155/2020/7846264>

СПИСОК ЛИТЕРАТУРЫ

1. Парфинович С.Н. Алгоритмы распознавания лиц для верификации личности по изображению. В сб.: *«Молодой исследователь: вызовы и перспективы»: сб. мат. CXIV Международной научно-практической конференции*. М.: Интернаука; 2019. С. 155–163.
2. Ахмедов А.А., Сагидов Г.С., Курбанисмаилов Г.М. Алгоритм распознавания лиц на основе метода Виолы – Джонса. В сб.: *«Молодой исследователь: вызовы и перспективы»: сб. мат. CXVIII Международной научно-практической конференции*. М.: Интернаука; 2019. С. 270–274.
3. Пентланд А., Чаудхари Т. Распознавание лиц для интеллектуальных сред. *Открытые системы*. 2000;03. URL: <https://www.osp.ru/os/2000/03/177939>
4. Горелик А.Л., Гуревич И.Б., Скрипкин В.А. *Современное состояние проблемы распознавания: Некоторые аспекты*. М.: Радио и связь; 1985. 161 с.
5. Самаль Д.И., Фролов И.И. Алгоритм подготовки обучающей выборки с использованием 3D-моделирования лиц. *Системный анализ и прикладная информатика*. 2016;4:17–23. URL: <https://sapi.bntu.by/jour/article/view/128/105>
6. Завалов Р.А., Гараев Р.А. Реализация алгоритма Виолы – Джонса на микроконтроллере с ограниченными ресурсами. *Наука и образование сегодня*. 2018;6(29):18–23. URL: <https://cyberleninka.ru/article/n/realizatsiya-algoritma-violy-dzhonsa-na-mikrokontrollere-s-ogranichennymi-resursami/viewer>
7. Балдин А.В., Елисеев Д.В. Алгебра многомерных матриц для обработки адаптируемой модели данных. *Наука и образование: научное издание МГТУ им. Н.Э. Баумана*. 2011;7:4. URL: <http://technomag.edu.ru/doc/199561.html>
8. Korotkov A. Database index for approximate string matching. In: *Proceedings of the 4th Spring/Summer Young Researchers' Colloquium on Software Engineering. SYRCoSE '10*. 2010, p. 136–140. <https://doi.org/10.15514/syrco-se-2010-4-27>
9. Кононыхин И.А., Ежов Ф.В., Мартынюк Р.А. и др. Реализация системы распознавания и отслеживания лиц. *Молодой ученый*. 2020;28(318):8–12. URL: <https://moluch.ru/archive/318/72492/>
10. Hinton G.E., Krizhevsky A., Wang S.D. Transforming auto-encoders. In: Honkela T., Duch W., Girolami M., Kaski S. (Eds.). *Artificial Neural Networks and Machine Learning – ICANN 2011. ICANN 2011. Lecture Notes in Computer Science*. Springer, Berlin, Heidelberg; 2011. V. 6791. P. 44–51. https://doi.org/10.1007/978-3-642-21735-7_6
11. Alghaili M., Li Z., Ali H.A.R. FaceFilter: Face identification with deep learning and filter algorithm. *Scientific Programming*. 2020;1–9. <https://doi.org/10.1155/2020/7846264>
12. Fitzgerald R.J., Price H.L., Valentine T. Eyewitness identification: Live, photo, and video lineups. *Psychology, Public Policy, and Law*. 2018;24(3):307–325. <http://dx.doi.org/10.1037/law0000164>
13. Etemad K., Chellapa R. Discriminant Analysis for Recognition of Human Face Images. *Journal of the Optical Society of America A*. 2004;14(8):1724–1733. <https://doi.org/10.1364/JOSAA.14.001724>

12. Fitzgerald R.J., Price H.L., Valentine T. Eyewitness identification: Live, photo, and video lineups. *Psychology, Public Policy, and Law*. 2018;24(3):307–325. <http://dx.doi.org/10.1037/law0000164>
13. Etemad K., Chellapa R. Discriminant analysis for recognition of human face images. *Journal of the Optical Society of America A*. 2004;14(8):1724–1733. <https://doi.org/10.1364/JOSAA.14.001724>
14. Kulikov A.A. The model is a reprint of an object in the image. *Rossiiskii tekhnologicheskii zhurnal = Russian technological journal*. 2020;8(3):7–13 (in Russ.). <https://doi.org/10.32362/2500-316X-2020-8-3-7-13>
15. Romanenko A.O., Yufryakov A.V. Image blur evaluation for biometric identification. *Nauka i obrazovanie segodnya = Science and Education Today*. 2018;7(30):16–19 (in Russ.). Available from URL: <https://cyberleninka.ru/article/n/otsenka-razmytiya-izobrazheniya-dlya-biometricheskoy-identifikatsii/viewer>
14. Куликов А.А. Модель репринта объекта на изображении. *Российский технологический журнал*. 2020;8(3):7–13. <https://doi.org/10.32362/2500-316X-2020-8-3-7-13>
15. Романенко А.О., Юфряков А.В. Оценка размытия изображения для биометрической идентификации. *Наука и образование сегодня*. 2018;7(30):16–19. URL: <https://cyberleninka.ru/article/n/otsenka-razmytiya-izobrazheniya-dlya-biometricheskoy-identifikatsii/viewer>

About the author

Alexander A. Kulikov, Cand. Sci. (Eng.), Associate Professor, Department of the Tool and Applied Software, Institute of Information Technologies, MIREA – Russian Technological University (78, Vernadskogo pr., Moscow, 119454 Russia). E-mail: tibult41@gmail.com. <https://orcid.org/0000-0002-8443-3684>

Об авторе

Куликов Александр Анатольевич, к.т.н., доцент, кафедра инструментального и прикладного программного обеспечения Института информационных технологий ФГБОУ ВО «МИРЭА – Российский технологический университет» (119454, Россия, Москва, пр-т Вернадского, д. 78). E-mail: tibult41@gmail.com. <https://orcid.org/0000-0002-8443-3684>

Translated by E. Shklovskii