

EQUALIZATION OF NONUNIFORM DISTRIBUTION OF RECOGNITION ERRORS PERCENTAGE OVER CLASSES IN CLASSIFYING SHIFTED MONOCHROME 60-BY-80-IMAGES

A study for equalizing nonuniform distribution of recognition errors percentage over classes is presented. The problem is exemplified on classifying shifted monochrome 60-by-80-images of 26 English alphabet letters. A general conception of adjusting the training process of two-layer perceptron classifier for the equalization is stated. In classifying shifted monochrome 60-by-80-images, this conception is used partially for cutting off extremely long training sample set. Within the example, the non-uniformity is reduced nearly for 25 %. The conception root is that the class representative recognized poorer is repeated in the training set. As the class recognition errors percentage increases, the repeat number is greater. However, the equalization conception is usable only for reasonable number of classes, so that the training sample set length could be shorter. Besides, the two-layer perceptron classifier can be adjusted by four parameters, determining the training sample set length and number of times when this set is passed through the perceptron. While equalizing, the hidden layer size of the two-layer perceptron should not be changed. For estimating non-uniformity, variance unbiased estimates may be used as well. The variance is decreased further using boosting ensembles of perceptron classifiers.

Keywords: object classification, shift, recognition errors percentage, distribution non-uniformity, perceptron, monochrome images, training sample.

B.B. ПОМАНИЮК
Хмельницький національний університет

ВИРІВНЮВАННЯ НЕРІВНОМІРНОГО РОЗПОДІЛУ ВІДСОТКОВОГО РІВНЯ ПОМИЛОК РОЗПІЗНАВАННЯ ЗА КЛАСАМИ У КЛАСИФІКАЦІЇ ЗСУНУТИХ МОНОХРОМНИХ ЗОБРАЖЕНЬ ФОРМАТУ 60-НА-80

Представляється дослідження з метою вирівнювати нерівномірний розподіл відсоткового рівня помилок розпізнавання за класами. Ця задача підкріплена прикладом класифікації зсунутих монохромних зображень 26 літер англійського алфавіту формату 60-на-80. Висвітлюється загальна концепція підлаштування навчального процесу класифікатора на основі двошарового перцептрону для описуваного вирівнювання. При класифікації зсунутих монохромних зображень формату 60-на-80 ця концепція використовується частково задля обрізання надзвичайно великої множини навчальної вибірки. У цьому прикладі нерівномірність знижується приблизно на 25 %. Суть висвітлюваної концепції полягає у тому, що представник певного класу, що розпізнається гірше, повторюється у навчальній множині. Зі зростанням відсоткового рівня помилок розпізнавання у цьому класі кількість повторів стає більшою. Однак дана концепція вирівнювання є застосовною лише для помірної кількості класів, щоб розмір множини навчальної вибірки був меншим. Крім того, класифікатор на основі двошарового перцептрону може бути підлаштований за чотирма параметрами, що визначають розмір множини навчальної вибірки та число разів, яке ця множина пропускається через перцептрон. Під час вирівнювання розмір прихованого шару двошарового перцептрону не змінюють. Для оцінювання нерівномірності також можуть бути застосовані й оцінки незміщеної дисперсії. Ця дисперсія знижується ще далі з використанням бустингових комітетів класифікаторів на основі перцептронів.

Ключові слова: класифікація об'єктів, зсув, відсотковий рівень помилок розпізнавання, нерівномірність розподілу, перцептрон, монохромні зображення, навчальна вибірка.

Problem of classification with nonuniform distribution of recognition errors percentage over classes

Performance of object classification mainly includes [1, 2] period of realtime classification, probability of the classification system failure, and recognition errors percentage (REP). Classification errors are plainly inevitable. Their estimations, obtained through testing the classification system, are needed for evaluating the classification quality on average [1, 3, 4]. An REP estimation $\tilde{r}_{ep}$ is compared to the maximally tolerated REP value $\hat{r}_{ep}$, and the classification system is implemented when $\tilde{r}_{ep} \leq \hat{r}_{ep}$. But here is a trivial example of that knowing REP is insufficient in controlling the classification system quality: for two classes, where $\tilde{r}_{ep} = 0.5 \cdot [r_{ep}(1) + r_{ep}(2)] = 9$ by $\hat{r}_{ep} = 10$, the first class REP is $r_{ep}(1) = 7$, and the second class REP is $r_{ep}(2) = 11$, whence it is clear that such classification quality is not admissible as system recognizes the second class too poor. Besides, even if REP in every class is less than $\hat{r}_{ep}$, but ratio of REP over some classes is far less or greater than one, then irregular distribution of REP over those classes is also unwanted. So, distribution of REP over classes is desirable to be uniform, and this is a problem for many fields of classification.

Assembling the classification system and controlling its quality

The contemporary classification system is usually a computer program or complex of computer programs, coded by algorithms on the basis of classification model. This model is a functional or an operator map, defined on the general totality of objects. The multilayer perceptron, being universal approximator, is one of the fastest classification models [2, 3, 5]. Its single hidden layer is sufficient for classifying almost everything. The two-layer perceptron classifier (TLPC) may be easily assembled in MATLAB, whereupon performance features of TLPC are being accumulated into MATLAB workspace and saved [1, 2, 4, 5]. Any performance feature, including REP over classes, TLPC acquires while being trained. Hence, control of TLPC quality is possible via adjusting the training process parameters. Particularly, it is about the training sample set feeding the input of TLPC.

Goal formulation

Let there be a general totality of objects from N classes, $N \in \mathbb{N} \setminus \{1\}$. Suppose that the trained TLPC performs with REP $r_{\text{ep}}(m)$ in m -th class, where the desired uniformity

$$r_{\text{ep}}(m) = N^{-1} \quad \forall m = \overline{1, N} \quad (1)$$

is not true, although

$$\frac{1}{N} \sum_{m=1}^N r_{\text{ep}}(m) = \tilde{r}_{\text{ep}} \leq \hat{r}_{\text{ep}}. \quad (2)$$

Moreover, the ultimately tolerated non-uniformity (UTNU), at which

$$\max \left\{ \frac{r_{\text{ep}}(m)}{r_{\text{ep}}(n)}, \frac{r_{\text{ep}}(n)}{r_{\text{ep}}(m)} : r_{\text{ep}}(n) > \tilde{r}_{\text{ep}}, m \in \{\overline{1, N}\}, n \in \{\overline{1, N}\} \right\} = r_{\text{max}} \leq r_1 \quad (3)$$

by an admissible REP $\tilde{r}_{\text{ep}}$ and some $r_1 > 1$, is not true also. That is there exist such $m_1 \in \{\overline{1, N}\}$ and $n_1 \in \{\overline{1, N}\}$ by $m_1 \neq n_1$ that

$$\max \left\{ r_{\text{ep}}(m_1) \cdot [r_{\text{ep}}(n_1)]^{-1}, r_{\text{ep}}(n_1) \cdot [r_{\text{ep}}(m_1)]^{-1} \right\} > r_1. \quad (4)$$

The goal is to adjust the training process of TLPC to get the condition (3) true. And the goal must be hit along with that the other performance features of TLPC should be changed minimally. This minimal-intruding theoretical adjustment for equalizing the nonuniform distribution of REP over classes should be demonstrated on a real example of classification.

Adjusting the training process of TLPC for REP equalization

A simple model of the object for its classification is monochrome image. For the monochrome image format $R \times C$ its model is $R \times C$ matrix of zeros and ones. So, the m -th class is represented as matrix $\mathbf{A}^{(m)} = (a_{rc}^{(m)})_{R \times C}$ at $a_{rc}^{(m)} \in \{0, 1\}$ by $r = \overline{1, R}$ and $c = \overline{1, C}$ for $m = \overline{1, N}$. While training, the input of TLPC is fed with the finite number of representatives of all the classes. Denote $\mathbf{A} = (\bar{a}_{jm})_{RC \times N}$ as the matrix of all original N monochrome $R \times C$ -images, reshaped into N columns. The input of TLPC is fed with U replicas of the matrix $\mathbf{A}$ along with F noised matrices $\{\mathbf{A}_k\}_{k=1}^F$, where the k -th noised matrix

$$\mathbf{A}_k = \Psi(\mathbf{A}, \mathbf{\Xi}_k) \quad (5)$$

is got after mapping the matrix $\mathbf{A}$ by the randomized parameters $\mathbf{\Xi}_k$. Particularly, the map Ψ can be just added for the noise as pixel distortion, where $\mathbf{\Xi}_k$ is usually $RC \times N$ matrix of values of normal variate [4, 6]. Generally, this map defines the character of distortions in original N images, because these distortions are model of differences between objects to be classified and representatives of N classes. The set

$$\tilde{\mathbf{P}}_{\text{train}}(NU + NF) = \{\tilde{\mathbf{P}}_i\}_{i=1}^{U+F} = \left\{ \{\mathbf{A}\}_{i=1}^U, \{\mathbf{A}_k\}_{k=1}^F \right\} \quad (6)$$

of original images and distorted images feeds the input of TLPC by the set of identifiers (targets)

$$\mathbf{T} = \{\mathbf{T}_i\}_{i=1}^{U+F} = \{\mathbf{I}\}_{i=1}^{U+F} \quad (7)$$

with identity $N \times N$ matrix $\mathbf{I}$, where number U indicates at how many replicas of undistorted images should be recognized in the training process. The set (6), being formed by (5), is passed through TLPC with identifiers (7) for Q_{pass} times. However, the trained TLPC with (6) by convention cannot ensure the inequality (3).

Let the map $\Omega \left(\left\{ r_{\text{ep}}^{(1)}(m) \right\}_{m=1}^N \right)$ transfer every normed REP

$$r_{\text{ep}}^{(1)}(m) = r_{\text{ep}}(m) \cdot \left(\sum_{n=1}^N r_{\text{ep}}(n) \right)^{-1} \quad (8)$$

into relative frequency $\lambda_{\text{rprs}}(m)$ of appearance (ARF) of the m -th class representatives in an updated set $\tilde{\mathbf{P}}_{\text{train}}$, $m = \overline{1, N}$ and $\sum_{m=1}^N \lambda_{\text{rprs}}(m) = 1$. In the simplest case, Ω is the identity map, and

$$r_{\text{ep}}^{(1)}(m) = \lambda_{\text{rprs}}(m) \quad \forall m = \overline{1, N}. \quad (9)$$

If the set $\{r_{\text{ep}}(m)\}_{m=1}^N$ were tolerated then the equality (9) would have meant that the greater REP in a class the more replicas of this class representatives must be included into the feeding set $\tilde{\mathbf{P}}_{\text{train}}$. Surely, sometimes this is far from actuality. Anyway, suppose that relative frequencies $\{\lambda_{\text{rprs}}(m)\}_{m=1}^N$ are known through the map of the normed

tolerated REP $\{r_{\text{ep}}^{(l)}(m)\}_{m=1}^N$ according to the inequality (3). Also let the operator $\mathfrak{F}$ be defined on the set of N terminating fractions with values from the half-interval $[0; 1)$ with their sum, being equal to 1. The operator $\mathfrak{F}$, having its argument as the set $\{\lambda_{\text{rprs}}(m)\}_{m=1}^N$, returns the denominator d_N , which is the least as possible after reducing every fraction from $\{\lambda_{\text{rprs}}(m)\}_{m=1}^N$ by the least common multiple. For instance, if

$$\{\lambda_{\text{rprs}}(m)\}_{m=1}^3 = \{0.3, 0.2, 0.5\}$$

then $\mathfrak{F}(\{0.3, 0.2, 0.5\}) = 10$, but for

$$\{\lambda_{\text{rprs}}(m)\}_{m=1}^5 = \{0.18, 0.2, 0.1, 0.22, 0.3\}$$

the denominator d_5 is

$$\mathfrak{F}(\{0.18, 0.2, 0.1, 0.22, 0.3\}) = 50.$$

Consequently, for

$$d_N = \mathfrak{F}\left(\{\lambda_{\text{rprs}}(m)\}_{m=1}^N\right) \quad (10)$$

there should be altogether $d_N \cdot (U + F)$ replicas of representatives of N classes in the updated set $\tilde{P}_{\text{train}}(d_N U + d_N F)$. This set feeds the input of TLPC with $d_N \cdot (U + F)$ replicas of representatives of N classes, what lets obtain the condition (3) for REP in all classes, unlike the set (6) with $N \cdot (U + F)$ replicas of representatives of N classes. Notably that the m -th class representative is repeated $U + F$ times in the set (6), and in the set $\tilde{P}_{\text{train}}(d_N U + d_N F)$ this representative is repeated $\lambda_{\text{rprs}}(m) \cdot d_N \cdot (U + F)$ times.

Thus adjusting the training process of TLPC for REP equalization consists in the new training set

$$\tilde{P}_{\text{train}}(d_N U + d_N F) = \left\{ \tilde{\mathbf{P}}_i \right\}_{i=1}^{U+F} = \left\{ \left\{ \tilde{\mathbf{A}}_l \right\}_{l=1}^U, \left\{ \tilde{\mathbf{A}}_k \right\}_{k=1}^F \right\} \quad (11)$$

that feeds the input of TLPC, wherein $\tilde{\mathbf{A}}$ and $\tilde{\mathbf{A}}_k$ are $RC \times d_N$ matrices, each of which contains $\lambda_{\text{rprs}}(m) \cdot d_N$ columns, representing the m -th class $\forall m = \overline{1, N}$. The set of targets is

$$T = \left\{ \tilde{\mathbf{T}} \right\}_{i=1}^{U+F}, \quad (12)$$

where $\tilde{\mathbf{T}} = [t_{mq}]_{N \times d_N}$ is matrix with $t_{mq} \in \{0, 1\}$, having $\lambda_{\text{rprs}}(m) \cdot d_N$ columns of their numbers within the set $C_N(m)$, where $t_{nq} = 1$ by $n = m$ and $t_{nq} = 0$ by $n \in \{\overline{1, N}\} \setminus \{m\}$ for $q \in C_N(m)$. Due to the operation (10) for columns of matrices in the set (11), the adjusted training process with the updated training set (11) for identifiers (12) must change performance features of TLPC minimally.

REP equalization for classification of monochrome 60-by-80-imaged alphabet letters

It is clear that for hitting the formulated goal explicitly there must be taken a certain general totality of objects. Therefore may 26 letters of English alphabet as monochrome images be the general totality. Here the letter is a generalized object model. At that may its format be 60×80 , not small, and not large, for efficient work with images.

Monochrome image distortion is commonly its shift. For treating this distortion type, some pixel color inversion is required during training TLPC [4]. Therefore let the map Ψ of the matrix $\mathbf{A} = (\bar{a}_{jm})_{4800 \times 26}$ shift (“shake”) every image from $\left\{ \mathbf{A}^{(m)} = (a_{rc}^{(m)})_{60 \times 80} \right\}_{m=1}^{26}$ in $\mathbf{A}$ with parameters Ξ_k , whereupon this column-shaken matrix is added to 4800×26 matrix of values of normal variate with zero expectation and some variance in Ξ_k , $k = \overline{1, F}$ (Figure 1 and Figure 2). In practice, it is sufficient to put $U = 2$ and $F = 8$ [4]. For a prompt exemplification, we take $U = F = 1$, and the added pixel distortion constitutes 10 % of the shift distortion. Further, the training set

$$\tilde{P}_{\text{train}}(52) = \{\tilde{\mathbf{P}}_1, \tilde{\mathbf{P}}_2\} = \{\mathbf{A}, \mathbf{A}_1\} \quad (13)$$

feeds the input of the MATLAB assembled TLPC by the set of targets

$$T = \{\mathbf{T}_1, \mathbf{T}_2\} = \{\mathbf{I}, \mathbf{I}\} \quad (14)$$

with identity 26×26 matrix $\mathbf{I}$, for $Q_{\text{pass}} = 50$ times. Unfortunately and expectedly, that the training process with the training set (13) by identifiers (14) makes TLPC perform with quite different REP over 26 classes (Figure 3).

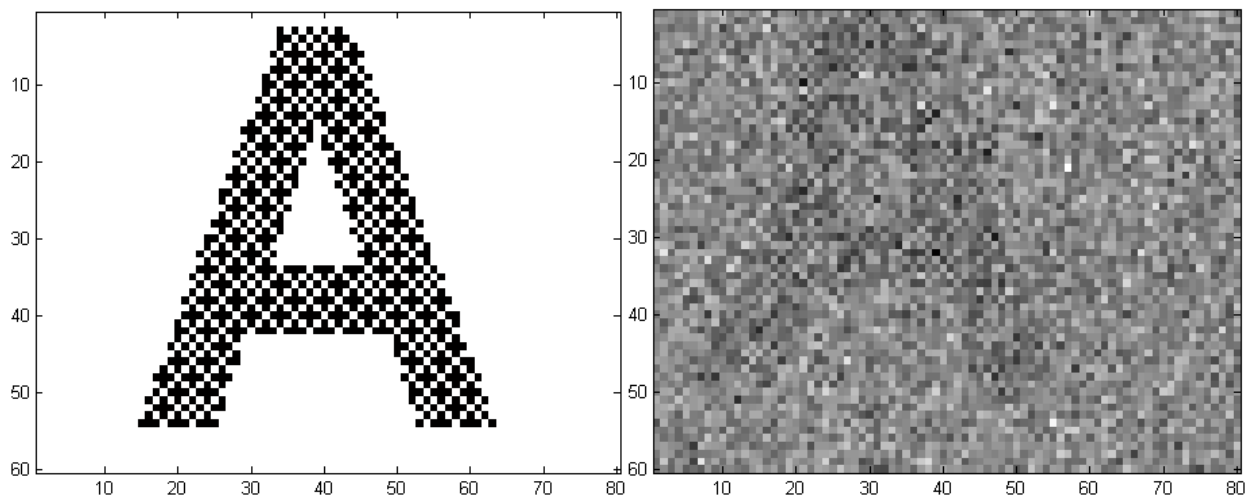


Fig. 1. Monochrome 60-by-80-image of letter “A” and its grayscale image after the “shaking” and pixel-distorting map (5), viewed within MATLAB



Fig. 2. Monochrome 60-by-80-image of letter “A” from Figure 1 after the “shaking” and pixel-distorting map (5), viewed within MATLAB

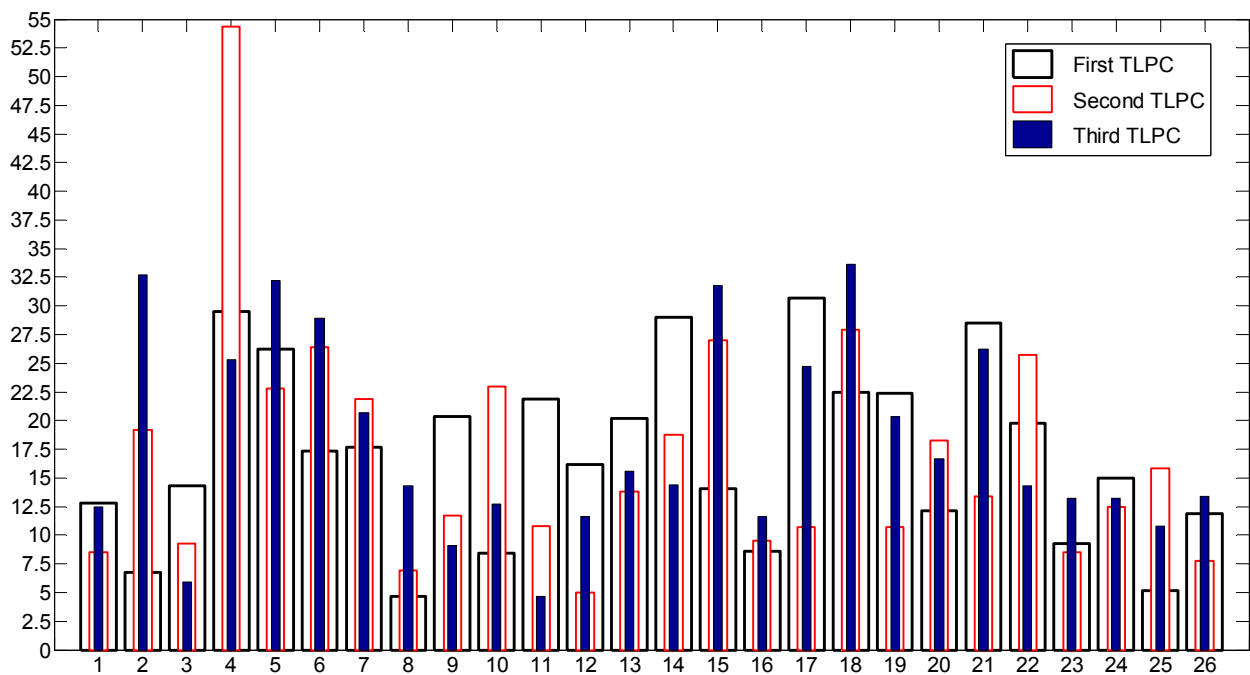


Fig. 3. Distributions of REP over 26 letters, derived from 1000 testings of three TLPC, trained with the training MATLAB function “traingda” at the fixed Ξ_1 for $U=1$, $F=1$, $Q_{\text{pass}}=50$ (the hidden layer size of TLPC is 250 neurons)

Statistically, these three REP estimations

$$\tilde{r}_{\text{ep}} \in \{17.1346, 16.9346, 18.0962\} \quad \text{by} \quad \tilde{r}_{\text{ep}} = \frac{1}{26} \sum_{n=1}^{26} r_{\text{ep}}(n) \quad (15)$$

appear less than the maximally tolerated REP value $\hat{r}_{\text{ep}} = 20$. Nevertheless, any of those distributions of $\{r_{\text{ep}}(m)\}_{m=1}^{26}$ over 26 classes has at least five letters recognized poorer (4, 5, 14, 17, 21 — for the first TLPC; 4, 6, 15, 18, 22 — for the second one; 2, 4, 5, 6, 15, 18, 21 — for the third one). In particular, letters “B”, “D”, “E”, “F”, “G”, “N”, “O”, “Q”, “R”, “U”, “V” are frequently classified wrong, whereas letter “H” is the most recognizable. The respective variance unbiased estimates [7, p. 212] are

$$v_{\text{ep}} \in \{60.1752, 106.6568, 74.3284\} \quad \text{by} \quad v_{\text{ep}} = \frac{1}{25} \sum_{m=1}^{26} \left(r_{\text{ep}}(m) - \frac{1}{26} \sum_{n=1}^{26} r_{\text{ep}}(n) \right)^2. \quad (16)$$

By the condition (3),

$$\begin{aligned} \max \left\{ r_{\text{ep}}(m) \cdot [r_{\text{ep}}(n)]^{-1}, r_{\text{ep}}(n) \cdot [r_{\text{ep}}(m)]^{-1} : r_{\text{ep}}(n) > 5, m \in \{1, 26\}, n \in \{1, 26\} \right\} = \\ = r_{\text{max}} \in \{5.9038, 7.8841, 5.6949\} \end{aligned} \quad (17)$$

for the tested three TLPC. Note that the second TLPC got so bad r_{max} because of the letter “D” training had failed.

To equalize REP distribution over 26 letters, there is the adjusting set (11)

$$\tilde{P}_{\text{train}}(2d_{26}) = \{\tilde{\mathbf{P}}_1, \tilde{\mathbf{P}}_2\} = \{\tilde{\mathbf{A}}, \tilde{\mathbf{A}}_1\} \quad (18)$$

by the denominator

$$d_{26} = \mathcal{F} \left(\{ \lambda_{\text{rprs}}(m) \}_{m=1}^{26} \right) \quad (19)$$

after having fulfilled the map

$$\Omega \left(\{ r_{\text{ep}}^{(1)}(m) \}_{m=1}^{26} \right) = \{ \lambda_{\text{rprs}}(m) \}_{m=1}^{26} \quad (20)$$

for the operator $\mathcal{F}$ argument in (19). The question of how to get the map in (20) has to be investigated within a separate paperwork. For now the very first suggestion is that this map is identity. Hence, in the updated training set (18) the m -th class representative is repeated $2d_{26}r_{\text{ep}}^{(1)}(m)$ times. That is the input of TLPC is fed with $4800 \times d_{26}$ matrices $\tilde{\mathbf{A}}$ and $\tilde{\mathbf{A}}_1$, each of which contains $d_{26}r_{\text{ep}}^{(1)}(m)$ columns, representing the m -th class $\forall m = \overline{1, 26}$. And the set $T = \{\tilde{\mathbf{T}}_1, \tilde{\mathbf{T}}_2\}$ of targets (12) consists of two replicas of matrix $\tilde{\mathbf{T}} = [t_{mq}]_{26 \times d_{26}}$ with $t_{mq} \in \{0, 1\}$, having $d_{26}r_{\text{ep}}^{(1)}(m)$ numbered within the set $C_{26}(m)$ columns, where $t_{nq} = 1$ by $n = m$ and $t_{nq} = 0$ by $n \in \{1, 26\} \setminus \{m\}$ for $q \in C_{26}(m)$.

The denominator (19) for each of those three tested TLPC should've been $d_{26} = 10000$. Obviously, this might have made the training sample extremely long, what would've have retarded the training process badly. Then we'll use a primitive routine for setting up d_{26} and $\{r_{\text{ep}}^{(1)}(m)\}_{m=1}^{26}$ or $\{\lambda_{\text{rprs}}(m)\}_{m=1}^{26}$. Say, the letters recognized poorer shall have a times greater ARF then the letters recognized better, $a > 1$. If there are N_1 letters recognized poorer and N_2 letters recognized better, $N_1 + N_2 = N$, then for ARF x of every letter recognized better we have:

$$N_1 a x + N_2 x = 1, \quad x = (N_1 a + N_2)^{-1} = (N + N_1(a - 1))^{-1}. \quad (21)$$

Taken $a = 2$ drives to that non-uniformity exposes even more, although the averaged REP decreases. Then, taking $a = 1.25$ for each of those three tested TLPC gives us the following. The first TLPC has

$$r_{\text{ep}}^{(1)}(m) = 5/109 \quad \text{by} \quad m \in \{4, 5, 14, 17, 21\} \quad \text{and} \quad r_{\text{ep}}^{(1)}(m) = 4/109 \quad \text{by} \quad m \in \{1, 26\} \setminus \{4, 5, 14, 17, 21\} \quad (22)$$

at $d_{26} = 109$; the second one has

$$r_{\text{ep}}^{(1)}(m) = 5/109 \quad \text{by} \quad m \in \{4, 6, 15, 18, 22\} \quad \text{and} \quad r_{\text{ep}}^{(1)}(m) = 4/109 \quad \text{by} \quad m \in \{1, 26\} \setminus \{4, 6, 15, 18, 22\} \quad (23)$$

at $d_{26} = 109$; the third TLPC has

$$r_{\text{ep}}^{(1)}(m) = 5/111 \quad \text{by} \quad m \in \{2, 4, 5, 6, 15, 18, 21\} \quad \text{and} \quad r_{\text{ep}}^{(1)}(m) = 4/111 \quad \text{by} \quad m \in \{1, 26\} \setminus \{2, 4, 5, 6, 15, 18, 21\} \quad (24)$$

at $d_{26} = 111$. Herewith, the training sample set length roughly quadruples. So, let $Q_{\text{pass}} = 10$. Now the training process with the training set (18) by $26 \times d_{26}$ identifiers within the set $T = \{\tilde{\mathbf{T}}_1, \tilde{\mathbf{T}}_2\}$ and the identity map (20) makes TLPC perform with smoothed distribution of REP over 26 letters (Figure 4), where each TLPC is tested twice. Besides, the REP estimation $\tilde{r}_{\text{ep}}$ has appeared being decreased significantly, and $\tilde{r}_{\text{ep}} < \hat{r}_{\text{ep}} = 13$:

$$r_{\text{ep}} \in \{12.2231, 12.2577\}, \quad r_{\text{ep}} \in \{10.7538, 10.2269\}, \quad r_{\text{ep}} \in \{11.4731, 11.1808\}. \quad (25)$$

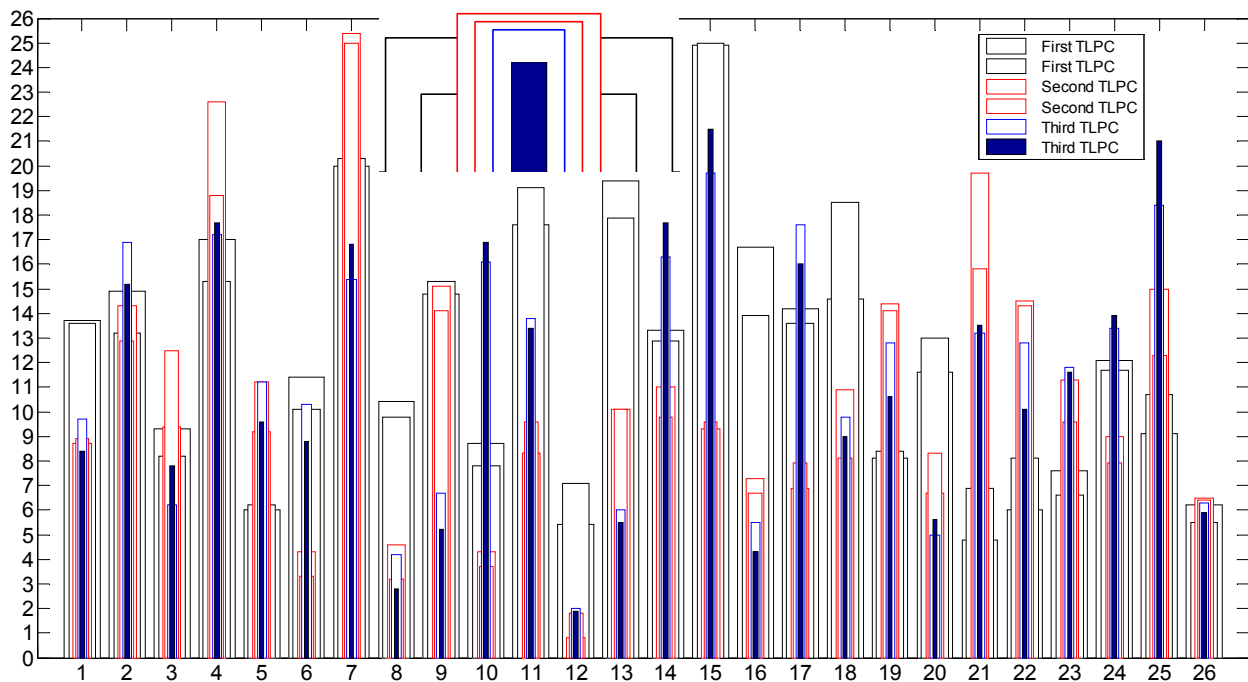


Fig. 4. Distributions of REP over 26 letters, derived from 1000 testings of TLPC after having trained it with the set (18) by the identity map (20) for MATLAB function “traingda” at the fixed Ξ_i for $U=1$, $F=1$, $Q_{\text{pass}}=10$ (the hidden layer size of TLPC is 250 neurons; each TLPC was initialized by the same randomizers used before to Figure 3)

A real proof that distributions of REP became smoother is that

$$v_{\text{ep}} \in \{26.4018, 24.9177\}, \quad v_{\text{ep}} \in \{32.5498, 25.8884\}, \quad v_{\text{ep}} \in \{25.0492, 30.9528\}. \quad (26)$$

By the condition (3),

$$r_{\text{max}} \in \{4.6111, 4.5455\}, \quad r_{\text{max}} \in \{3.9077, 3.9062\}, \quad r_{\text{max}} \in \{3.5818, 4.1346\}. \quad (27)$$

And comparison of (27) to (17) tells that, circumspetively, the non-uniformity has been reduced almost for one fourth. Notwithstanding this, letters “G”, “K”, “Y” became less recognizable on average.

Noteworthy, that the set (18), being passed through TLPC with identifiers $T = \{\tilde{T}_1, \tilde{T}_2\}$ for $Q_{\text{pass}} = 10$ times, lingers the training process in comparing it to passing the set (13) through TLPC with identifiers (14) for the same parameters. The relative delay of the training process can be considered as ratio of traintime for the set (18) to traintime for the set (13). For the demonstrated example this ratio is about 3:1, although number of columns in the set (18) is almost four times greater than number of columns in the set (13). Nonetheless the efficiency of REP equalization for classification of monochrome 60-by-80-imaged alphabet letters can be increased by optimizing the hidden layer size of TLPC, what may shorten the elongate traintime.

Conclusion

Equalization of nonuniform distribution of REP over classes requires to extend the training set for TLPC, including into this set the repeated class representatives. Number of the class representative replicas is proportional to product of the map of the normed REP in this class into this class representative ARF and the least as possible denominator for ARF in all classes. Namely, within the extended training set for TLPC the m -th class representative replicas number is $\lambda_{\text{rps}}(m)d_N$ times greater than for the initial training set.

However, there are a lot of difficulties in accomplishing the equalization. Firstly, determination of the map Ω is a problem, whose solution with assuming this map to be the identity is sometimes poor. Secondly, before determining the map Ω the distribution of REP $\{r_{\text{ep}}(m)\}_{m=1}^N$ must be very stable for the ascertained general totality of objects. This is necessary because every next time TLPC is retrained with the training set (6), there are statistically unstable distributions of REP for TLPC, especially if it is tested insufficiently. In REP equalization for classification of monochrome 60-by-80-imaged alphabet letters, it needs hundreds of cycles of roughly 1000 testings for obtaining the stable averaged distribution of $\{r_{\text{ep}}(m)\}_{m=1}^{26}$ over 26 letters. As an exclusion, TLPC for a specific classification problem, trained with the training set (6), may be retrained further with the new training set (11), though such two-staged training process will make TLPC effectiveness reliable only locally. Another faster way to try equalizing distribution of REP for TLPC is to spot some classes with the highest REP and add replicas of their representatives into an updated training set. But the adding shouldn't be much as else the updated distribution of REP becomes yet more unsmooth than previously. In this way, the example with $a=1.25$ and (21) for classifying shifted monochrome 60-by-80-images came out half-successful.

The third point, which has to be mentioned necessarily, is the case when a part of ARF set $\{\lambda_{\text{rps}}(m)\}_{m=1}^N$ consists of small numbers, might be rounded to zero. This may occur especially for great N . That causes to use greater d_N what lingers the training process. Henceforward, regarding the spoken three factors, if UTNU condition (3) is not true, but the condition (2) is true for a pretty low-valued $\hat{r}_{\text{ep}}$ then the process of equalization of nonuniform distribution of REP over classes better should not be started. And if there is inequality (4) along with $r_{\text{ep}}(m_1) > \hat{r}_{\text{ep}}$ or $r_{\text{ep}}(n_1) > \hat{r}_{\text{ep}}$ then REP over m_1 -th or n_1 -th class correspondingly must be corrected anyway.

Finally, there is a list of adjustable parameters for TLPC to be trained for REP equalization by the training set (11). They are a , U , F , and a new Q_{pass} . It is unlikely that $a \geq 2$ gives smoother REP distribution. So, $a \in (1; 2)$ or better try $a \in (1; 1.5)$ instead. Integers U and F mustn't be increased. The new Q_{pass} is recommended to be b times lesser than the old Q_{pass} for the training set (6), where b is taken approximately equal to $\frac{d_N}{N}$. The hidden layer size of TLPC should not be changed.

When non-uniformity of distribution of REP over classes is estimated, see also variance unbiased estimates

$$v_{\text{ep}} = \frac{1}{N-1} \sum_{m=1}^N \left(r_{\text{ep}}(m) - \frac{1}{N} \sum_{n=1}^N r_{\text{ep}}(n) \right)^2 \quad (28)$$

along with r_{max} in UTNU condition (3). Both variance estimate and r_{max} must decrease. There is no obvious priority. A combination of the normalized estimate (28) and r_{max} is possible. This combination will answer the question of how good the REP equalization is. Once the variance (28) is decreased for a set of TLPC, it is expedient to use these TLPC in ensemble for boosting, which decreases the variance (28) further [5] and assists in REP equalization additionally.

References

1. Romanuke V. V. Model of optimizing the multistage process of neuronet training with wastage constraint function in an image recognition problem, Herald of Vinnytsya polytechnic institute, 2013, N. 1, pp. 104—109.
2. Romanuke V. V. Combining the backpropagation algorithm training functions for a bit map character recognition problem under noise with feed-forward neural network, Scientific bulletin of Chernivtsi National University of Yuriy Fedkovych. Series: Computer systems and components, 2012, Vol. 3, Iss. 1, pp. 75—80.
3. Arulampalam G., Bouzerdoum A. A generalized feedforward neural network architecture for classification and regression, Neural Networks, 2003, Vol. 16, Iss. 5—6, pp. 561—568.
4. Romanuke V. V. An attempt for 2-layer perceptron high performance in classifying shifted monochrome 60-by-80-images via training with pixel-distorted shifted images on the pattern of 26 alphabet letters, Radioelectronics, informatics, control, 2013, N. 2, pp. 112—118.
5. Romanuke V. V. Accuracy improvement in wear state discontinuous tracking model regarding statistical data inaccuracies and shifts with boosting mini-ensemble of two-layer perceptrons, Problems of tribology, 2014, N. 4, pp. 55—58.
6. Romanuke V. V. Setting the hidden layer neuron number in feedforward neural network for an image recognition problem under Gaussian noise of distortion, Computer and Information Science, 2013, Vol. 6, N. 2, pp. 38—54.
7. Gmurman V. Y. Probability theory and mathematical statistics, 9-th ed., Moscow, Vyssh. sc., 2003, 479 p.

Література

1. Романюк В. В. Модель оптимізації багатоетапного процесу нейромережевого навчання з функцією обмеження втрат у задачі розпізнавання зображень / В. В. Романюк // Вісник Вінницького політехнічного інституту. — 2013. — № 1. — С. 104 — 109.
2. Romanuke V. V. Combining the backpropagation algorithm training functions for a bit map character recognition problem under noise with feed-forward neural network / V. V. Romanuke // Науковий вісник Чернівецького національного університету імені Юрія Федьковича. Серія: Комп'ютерні системи та компоненти. — Чернівці : ЧНУ, 2012. — Т. 3, вип. 1. — С. 75 — 80.
3. Arulampalam G. A generalized feedforward neural network architecture for classification and regression / G. Arulampalam, A. Bouzerdoum // Neural Networks. — 2003. — Volume 16, Issues 5 — 6. — P. 561 — 568.
4. Romanuke V. V. An attempt for 2-layer perceptron high performance in classifying shifted monochrome 60-by-80-images via training with pixel-distorted shifted images on the pattern of 26 alphabet letters / V. V. Romanuke // Радіоелектроніка, інформатика, управління. — 2013. — № 2. — С. 112 — 118.
5. Romanuke V. V. Accuracy improvement in wear state discontinuous tracking model regarding statistical data inaccuracies and shifts with boosting mini-ensemble of two-layer perceptrons / V. V. Romanuke // Problems of tribology. — 2014. — N. 4. — P. 55 — 58.
6. Romanuke V. V. Setting the hidden layer neuron number in feedforward neural network for an image recognition problem under Gaussian noise of distortion / V. V. Romanuke // Computer and Information Science. — 2013. — Vol. 6, N. 2. — P. 38 — 54.
7. Гмурман В. Е. Теория вероятностей и математическая статистика : [учеб. пособие для вузов] / Гмурман В. Е. — [9-е изд., стер.]. — М. : Высш. шк., 2003. — 479 с. : ил.

Рецензія/Peer review : 20.3.2015 р.

Надрукована/Printed : 15.4.2015 р.

Рецензент: стаття рецензована редакційною колегією